

User Manual and Technical Guidance for The Bayesian Benchmark Dose (BBMD) Analysis System

Version 3.0.2 (2025.7.30)

This document primarily serves as a user manual to provide introductions for using the BBMD system properly. Technical details are also provided to help users better understand the Bayesian statistical methodology and other algorithms employed in this computational system and interpret the results.

I.	OVERVIEW OF THE BBMD SYSTEM	4
A.	How to Access the System	4
B.	Create and Log in to an account	5
C.	The User Dashboard	5
i.	<i>Access Previous Analyses</i>	<i>6</i>
a.	How to review a previous analysis	6
b.	How to edit a previous analysis	7
c.	How to share a previous analysis	7
ii.	<i>Add New Analyses</i>	<i>8</i>
II.	BMD ANALYSIS FOR A SINGLE DATASET	11
A.	General Introduction on BMD Analysis	11
B.	Data Input and Pre-analysis	11
C.	MCMC Settings	15
i.	<i>How to make and change settings</i>	<i>15</i>
ii.	<i>How MCMC settings may impact the results</i>	<i>16</i>
D.	Model settings	16
i.	<i>Dose-Response Models for Dichotomous Data</i>	<i>19</i>
ii.	<i>Dose-Response Models for Continuous Data</i>	<i>20</i>
iii.	<i>Dose-Response Models for Categorical Data</i>	<i>23</i>
E.	Model Fit Results	24
i.	<i>Parameter estimation results</i>	<i>26</i>
ii.	<i>Posterior Predictive P-Value</i>	<i>26</i>
iii.	<i>Posterior Model Weight</i>	<i>27</i>
iv.	<i>Interactive Dose-Response Plot</i>	<i>28</i>
v.	<i>Correlation Matrix</i>	<i>28</i>
vi.	<i>Plots for parameter posterior sample</i>	<i>28</i>
F.	BMD Estimation	28
i.	<i>Dichotomous Data</i>	<i>28</i>
a.	How to execute a BMD Analysis	28
b.	Explanation of the analysis calculation	31
ii.	<i>Continuous Data</i>	<i>31</i>
a.	How to execute a BMD Analysis	31
b.	Explanation of the analysis calculation	34
iii.	<i>Categorical Data</i>	<i>35</i>
a.	How to execute a BMD Analysis	35
b.	Explanation of the analysis calculation	36
iv.	<i>Model Averaged BMD Calculation</i>	<i>36</i>
a.	Posterior Model Weight for Model Averaged BMD Calculation	37
b.	Model Averaged BMD Calculation	37
G.	Probabilistic RfD Estimation	37
H.	Risk/Response at Dose	40
I.	Share the Analysis	42
III.	BATCH PROCESSING FOR ANALYSIS	43

A.	Data Input	43
B.	Model Settings	43
C.	MCMC Settings	44
D.	BMR Settings	45
E.	Execute	45
F.	Results	46
G.	Share/Review	46
IV.	BMD ANALYSIS FOR GENOMIC DATA	48
A.	Introduction to BMD analysis for Genomic Data	48
B.	Data input	48
i.	<i>How to input data into the system</i>	48
ii.	<i>How to interpret the data input results</i>	49
C.	Data Pre-processing	50
i.	<i>How to perform data pre-processing</i>	50
ii.	<i>How to interpret the data pre-processing results</i>	51
iii.	<i>Explaining the algorithms used</i>	52
a.	Fold change	52
b.	P-value and adjusted P-value	53
D.	Data Preparation	53
E.	Model Selection	54
F.	BMD Settings	56
i.	<i>Genomic BMD analysis steps</i>	56
ii.	<i>Explanation of the analysis calculation</i>	57
G.	MCMC Settings	57
H.	BMD Results	58
I.	Platform Selection	59
J.	Pathway Analysis Results	60
i.	<i>How to view the pathway analysis results</i>	60
ii.	<i>Algorithms used for pathway analysis</i>	61
K.	Share/Review	62
V.	PROBABILISTIC RfD ANALYSIS	64
A.	Introduction to the Probabilistic RfD Analysis	64
B.	Perform a Probabilistic RfD Analysis	64
C.	Interpret Probabilistic RfD Analysis Results	67
VI.	BMD ANALYSIS FOR EPIDEMIOLOGICAL DATA	68
A.	Introduction to BMD analysis for Epidemiological Data	68
B.	Data input	68
i.	<i>Choosing an exposure type</i>	68
ii.	<i>How to input data into the system</i>	68
C.	MCMC Settings	69
i.	<i>How to make and change settings</i>	69
ii.	<i>How MCMC settings may impact the results</i>	70
D.	Model Settings	70
E.	Model Fit Results	70
i.	<i>Parameter estimation results</i>	72
ii.	<i>Posterior Predictive P-Value</i>	72
iii.	<i>Posterior Model Weight</i>	73
iv.	<i>Interactive Dose-Response Plot</i>	73
v.	<i>Correlation Matrix</i>	74
vi.	<i>Plots for parameter posterior sample</i>	74
F.	BMD Estimation	74
VII.	MULTISITE TUMOR BMD ANALYSIS	77

A.	Introduction to Multisite Tumor BMD Analysis	77
B.	Data input	77
i.	<i>Setting your tumor sites</i>	77
ii.	<i>How to input data into the system</i>	77
C.	MCMC Settings	78
i.	<i>How to make and change settings</i>	78
ii.	<i>How MCMC settings may impact the results</i>	79
D.	Model Settings	79
E.	Model Fit Results	80
i.	<i>Parameter estimation results</i>	81
ii.	<i>Posterior Model Weight</i>	81
iii.	<i>Individual Tumor Site Models</i>	82
F.	BMD Estimation	82
VIII.	REFERENCES	83

I. Overview of the BBMD System

Welcome to BBMD, the Bayesian Benchmark Dose Modeling system! BBMD is a robust tool for probabilistic dose-response assessment. Five types of dose-response analysis are available including (a) BMD analysis for single dataset, (b) batch processing for BMD analysis, (c) BMD analysis for genomic data, (d) probabilistic reference dose analysis, (e) BMD analysis for epidemiological data, and (f) Specialty models for multisite tumors and nested dichotomous (coming soon). Each type of analysis is discussed in detail from Sections 2-7. This chapter contains a description of how to assess BBMD, and an overview of the BBMD system.

A. How to Access the System

The URL of the BBMD online system is [Bayesian BMD \(benchmarkdose.com\)](https://benchmarkdose.com). Chrome and Firefox are the two recommended web-browsers for using the system. The front page of the BBMD system is displayed in Figure 1.1 below.

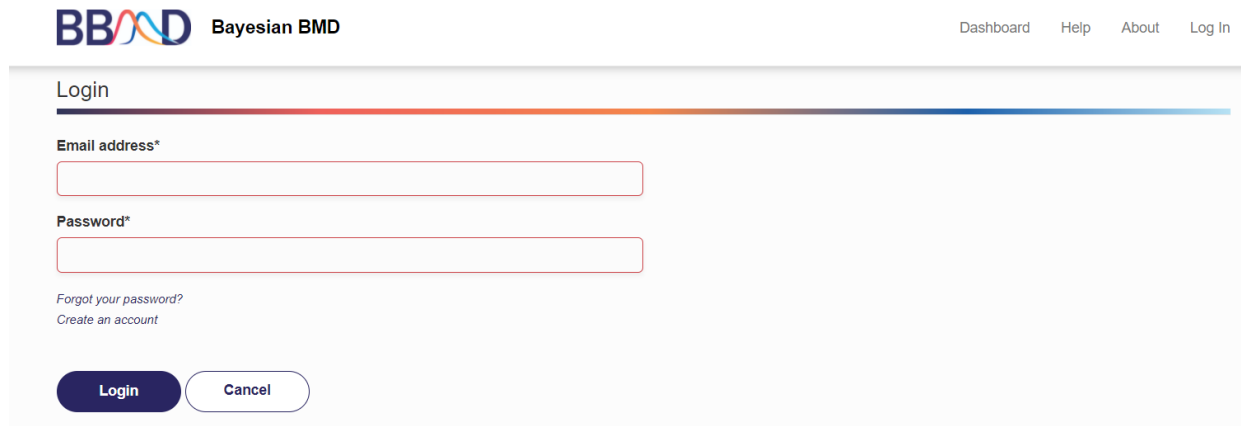


Figure 1.1. Front page of the BBMD system

In the upper right-hand corner of the webpage (Figure 1.1), there are five different options. These five options are available on every page of the system. The first option, “Dashboard”, which is where you can access your analyses. This includes analyses you have created and analyses that have been shared with you, which are stored in their respective tabs. The next option in the right corner is “Help”. This link takes you to the download of this user manual, something you must have already figured out. The third option, “About”, is a brief overview on the BBMD system, which gives you 1) a summary of this system, 2) references on the methodology of BBMD, 3) preferred citation of this website, 4) contact info of the DREAM Tech LLC development team for any questions or suggestions, and finally 5) the funding agency. The fourth option, “FAQ”, listed frequently asked questions for BBMD system. The final option in the upper right-hand corner is “Log in”. To use the system for BMD analysis, users first need to either log in to the system.

B. Create and Log in to an account

By clicking the “Log-in” option, the “Log-in” page is shown in Figure 1.2 below. An email address (i.e., the username) and your password are required when logging in. If this is your first time using this system, you will need to create an account by clicking the “Create an account” link at the bottom of the login page shown in Figure 1.2.

The screenshot shows the BBMD (Bayesian BMD) login page. At the top left is the BBMD logo and the text "Bayesian BMD". At the top right are links for "Dashboard", "Help", "About", and "Log In". The main heading is "Login". Below it are two input fields: "Email address*" and "Password*". Below the password field are two links: "Forgot your password?" and "Create an account". At the bottom are two buttons: "Login" and "Cancel".

BBMD Bayesian BMD

Dashboard Help About Log In

Login

Email address*

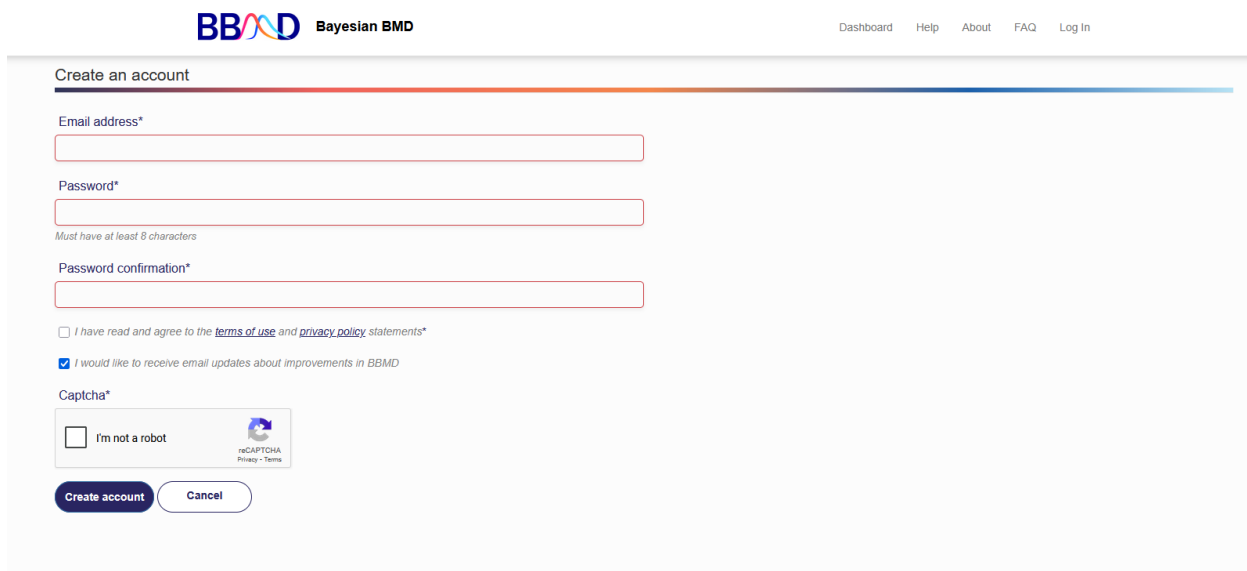
Password*

[Forgot your password?](#)
[Create an account](#)

Login Cancel

Figure 1.2. Login page of the BBMD system

The “Create an account” page is shown in the Figure 1.3. Just like most online systems, an email address (used as an account name) and password, which is entered twice, are needed to create an account. By creating a personal account, your previous analyses will be saved in your account for future review or update. A confirmation email will be sent to your inbox for verification after creating an account. Click on the link in the email to begin using the BBMD software.

The screenshot shows the BBMD (Bayesian BMD) "Create an account" page. At the top left is the BBMD logo and the text "Bayesian BMD". At the top right are links for "Dashboard", "Help", "About", "FAQ", and "Log In". The main heading is "Create an account". Below it are three input fields: "Email address*", "Password*", and "Password confirmation*". Below the password fields is a note: "Must have at least 8 characters". Below the confirmation field are two checkboxes: "I have read and agree to the terms of use and privacy policy statements*" (unchecked) and "I would like to receive email updates about improvements in BBMD" (checked). Below these is a "Captcha*" section with an "I'm not a robot" checkbox and a reCAPTCHA logo. At the bottom are two buttons: "Create account" and "Cancel".

BBMD Bayesian BMD

Dashboard Help About FAQ Log In

Create an account

Email address*

Password*

Must have at least 8 characters

Password confirmation*

☐ I have read and agree to the [terms of use](#) and [privacy policy](#) statements*

☒ I would like to receive email updates about improvements in BBMD

Captcha*

☐ I'm not a robot

reCAPTCHA Privacy - Terms

Create account Cancel

Figure 1.3. “Create an account” Page of the BBMD System

C. The User Dashboard

i. Access Previous Analyses

a. How to review a previous analysis

Once logged in, you will reach the page with options to access your previous analyses (if any are completed), or analyses shared with you (if you have any) or analyses you shared with the experts from Dream Tech LLC (if you have any) in the respective tabs. If you select one of your previous analyses (e.g., “EXAMPLE” in Figure 1.4), the previous results will be shown as the one is shown in Figure 1.5, you can review the previous analysis. For each type of analysis, there are unique tabs available. By clicking each of the tabs, you can access the data or settings stored in this analysis. One or more tabs may be empty, which indicates that the analysis was previously ended before the corresponding section had been finished. The detailed information regarding the contents in each of the tabs will be introduced in the following sections.

Hello, Welcome to your DREAM Dashboard

My Analyses

Shared with Me

Dream Assist

+ New analysis

Name	Data type	Summary	Date created	Last updated	Status
Single Dataset Runs					
EXAMPLE	Dichotomous (summary)	8 Model(s), 1 BMD(s), 0 RAD(s)	02/13/2024 2:31 PM	03/26/2025 9:27 AM	Shared
Video Example 2	Dichotomous (summary)	8 Model(s), 0 BMD(s), 0 RAD(s)	02/13/2024 2:36 PM	02/13/2024 2:37 PM	Private
Workshop Public Example	Dichotomous (summary)	8 Model(s), 0 BMD(s), 0 RAD(s)	01/17/2023 3:36 PM	01/17/2023 3:37 PM	Private
New run Dec 09 2022, 01:48 PM	Continuous (summary)	8 Model(s), 0 BMD(s), 0 RAD(s)	12/09/2022 1:48 PM	12/09/2022 1:48 PM	Private
New run May 06 2022, 07:10 PM	Continuous (summary)	0 Model(s), 0 BMD(s), 0 RAD(s)	05/06/2022 7:10 PM	05/06/2022 7:10 PM	Private
New run May 06 2022, 07:08 PM	Continuous (summary)	0 Model(s), 0 BMD(s), 0 RAD(s)	05/06/2022 7:08 PM	05/06/2022 7:08 PM	Private
SOT Categorical Example	Categorical	3 Model(s), 1 BMD(s), 1 RAD(s)	03/24/2022 10:04 AM	03/24/2022 10:05 AM	Private
SOT Dichotomous Example	Dichotomous (summary)	8 Model(s), 1 BMD(s), 1 RAD(s)	03/24/2022 9:59 AM	03/24/2022 9:59 AM	Private
SOT Continuous Example	Continuous (individual)	8 Model(s), 1 BMD(s), 1 RAD(s)	03/24/2022 9:52 AM	03/24/2022 9:53 AM	Private
Batch Analyses					
Genomic Analyses					

Figure 1.4. The Summary Page of Existing Runs

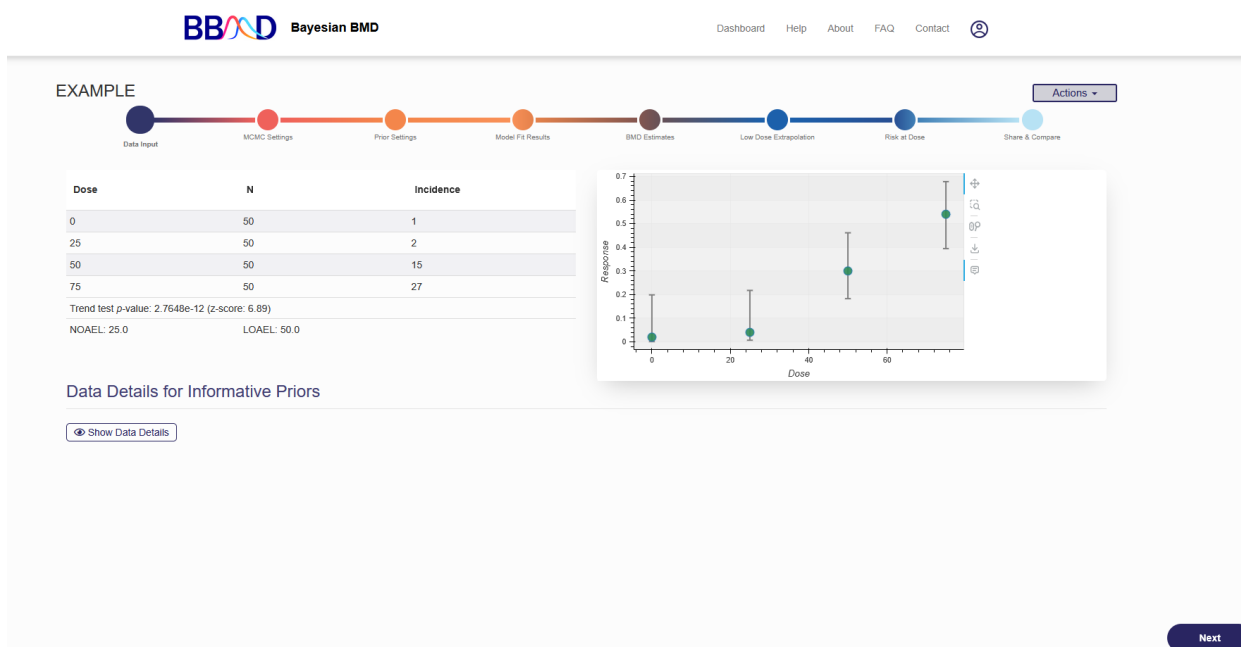


Figure 1.5. Reviewing an Existing Analysis

b. How to edit a previous analysis

You can also edit previous analysis. The pull-down menu of “Action” in Figure 1.6 has the “Update”, “Delete”, and “Edit a Copy” options. “Update” takes you to the editing phase of an analysis session where you can specify analysis settings and execute analysis estimations. Selecting “Delete” will prompt you to confirm you want to delete this analysis. If you choose to delete an analysis, the analysis and all results will be removed from your profile. Selecting “Edit a Copy” will create an editable copy of a previous analysis.

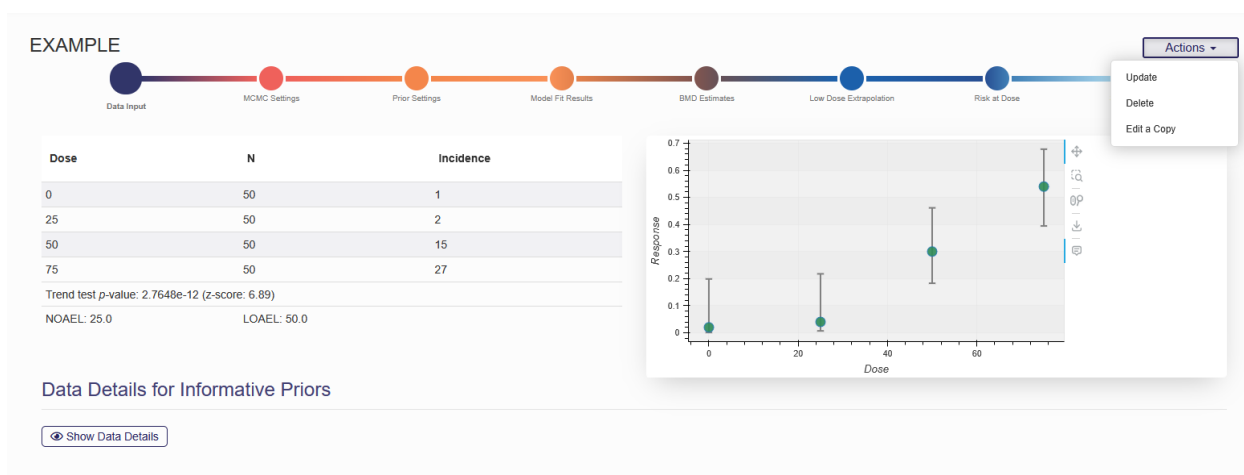


Figure 1.6. The Actions Available to Users for an Existing Analysis

c. How to share a previous analysis

The final tab in every type of analysis is the “Share” tab, shown in Figure 1.7. From this page you can send your analysis to yourself, share your analysis with the public, allow other people to edit

the analysis, and even send your analysis to the DREAM Tech team who can assist with the analysis. If wish to share your analysis with the public, you will be given a share link which can be sent to different users. You can also invite users directly using their account email.

Under the “share” tab, you may also be able to export the results of the analysis into Word or Excel formats, depending on the type of analysis. Exporting the results will send a link to your account’s email where you can download the reports.

At the bottom, you can compare the BMD estimates generated by the BBMD system with those from the EPA’s BMD calculations. Please note you must have at least 1 BMR definition. By clicking the button “Execute”, your data will be sent to the EPA's BMDS online server for executing.

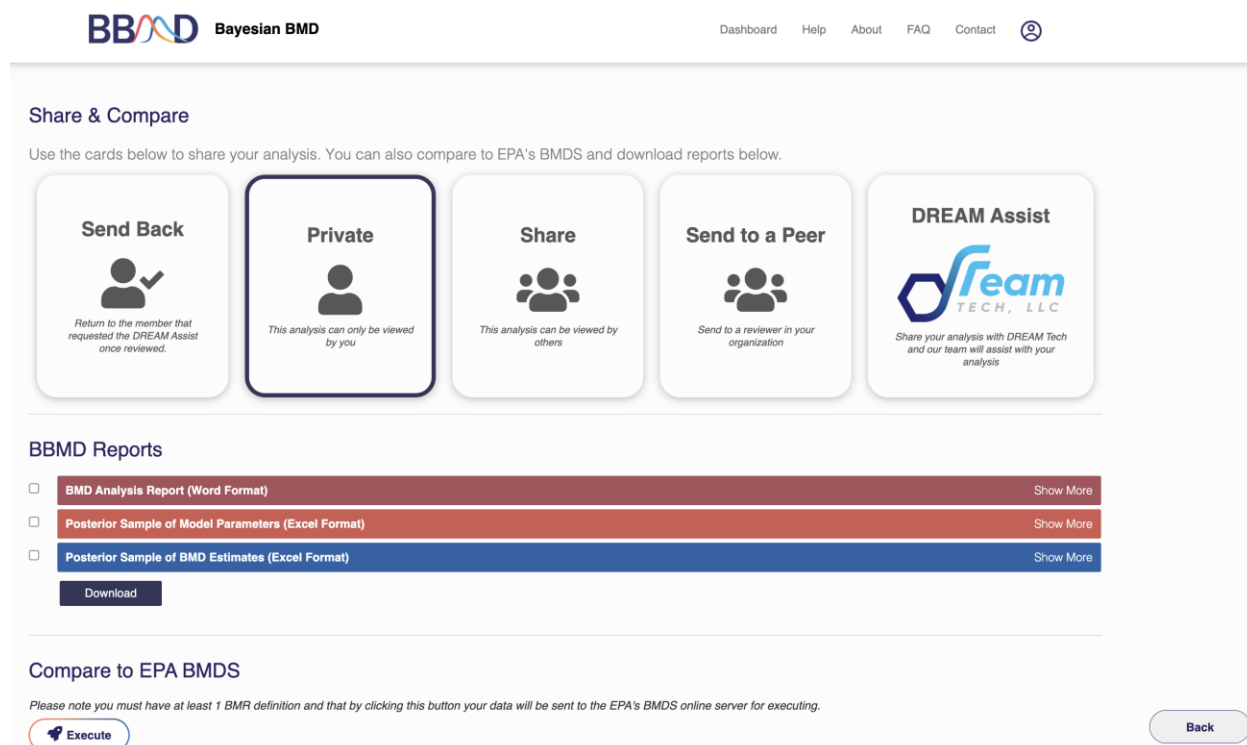


Figure 1.7. Sharing an analysis

ii. Add New Analyses

To start a new analysis from scratch, click “New analysis” in the top right corner in Figure 1.4. When starting a new analysis, you will be asked to select which type of analysis you would like to perform (Figure 1.8). The current analysis types include:

- BMD Analysis for Single Dataset,
- Batch Processing for BMD Analysis,
- BMD Analysis for Genomic Data,
- Probabilistic Reference Dose (RfD) Analysis,
- BMD Analysis for Epidemiological Data,
- Specialty models for multisite tumors and nested dichotomous (coming soon).

After specifying the data type on the analysis selection page, you will be directed to the corresponding module. If the “BMD Analysis for Single Dataset” module or “Batch Processing for BMD Analysis” module is chosen, the data type either continuous or dichotomous or categorical also needs to be specified. If the “Specialty Models” module is chosen, the data type either multisite tumors or nested dichotomous (coming soon) also needs to be specified. These data types are described below. For each type of analysis, a detailed explanation of the required inputs, modeling settings, and model outputs are in sections 2-7.

- Continuous – A continuous response is reported as a measurement of the effect, such as body or tissue weights, in control and exposure groups. The response may be reported in either absolute or relative change from control. When individual data are available, the dose and response data input to BBMD for BMD inference. Instead, when individual data are not available, the summary data are needed including dose level, number of subjects in each dose group, mean value of the response, and the standard deviation or standard error of the response.
- Dichotomous - A dichotomous response is reported as either the presence or absence of an effect. Dichotomous data are reported either as the summary data with the number of animals showing the effect at each individual level or as the individual data with “0” or “1” indicating that the subject is non-affected or affected respectively at each dose level. In BBMD, dichotomous summary data require three values for each dose group (i.e., each input row): dose level, total number of subjects in that group, and the number of subjects affected. Dichotomous individual data require two values for each input row (representing each subject): dose level, and “0” or “1” indicating that the subject is non-affected or affected respectively.
- Categorical – A categorical response is classified as the one or more defined category (e.g., mild, moderate, or severe change) in addition to the no-effect category. This type requires three values for each input row representing each individual subject: dose, severity level, and response.

Please Select Analysis Type

[← Go Back to Dashboard](#)**BMD Analysis for Single Dataset**

This analysis is designed to provide detailed results for a single data set of various types, including model fitting results, response at a dose estimation, BMD estimation, probabilistic RfD estimation, etc. Please Select Data Type

Data Types

[Continuous](#) [Dichotomous](#) [Categorical](#)**Batch Processing for BMD Analysis**

This analysis is designed to provide an efficient but slightly simplified way to analyze a large number of datasets all together. Please Select Data Type

Data Types

[Continuous](#) [Dichotomous](#) [Categorical](#)**BMD Analysis for Genomic Data**

This analysis is designed to perform BMD modeling and estimation for genomic data.

Probabilistic Reference Dose (RfD) Analysis

This analysis can be used to convert traditionally estimated BMD/BMDL, NOAEL/LOAEL to probabilistic reference dose given some additional input information.

BMD Analysis for Epidemiological Data

This analysis is especially designed for epidemiological dose-response data published in literature, i.e., subjects are grouped according to exposure ranges. If raw individual subject exposure and response data are available, you may consider using the modeling method under BMD Analysis for Single Dataset.

Specialty Models

These models are designed for specialty data types. **Nested Dichotomous models are in development and will be released soon.**
Please Select a Specialty Model

Data Types

[Multisite Tumor](#) [Nested Dichotomous](#)[Create](#)

Figure 1.8. Analysis type selection screen when beginning a new analysis

II. BMD Analysis for a Single Dataset

A. General Introduction on BMD Analysis

After you selected “BMD Analysis for Single Dataset” and specified the data type in Figure 1.8, the web page will change to the one as shown in Figure 2.1. An automatically generated name, “New Continuous/Dichotomous/Categorical Run *Month Day Year, HH:MM AM/PM*”, is assigned to the newly started analysis. You can click the pencil button next to the analysis name, as shown in the Figure 2.1, to make the name more identifiable. Without inputting any data, you will not be able to continue to advance through the tabs. The first step in an analysis is to input dose-response data. For continuous and dichotomous data, whether the dataset is “summary” or “individual” needs to be specified. After each step is finished, the “Next” button and the tab in the progress bar will light up.

The screenshot shows the 'Manual Example' page of the Bayesian BMD web application. At the top, a progress bar contains ten steps: Data Input, MCMC Settings, Model Settings, Prior Settings, Execute Model Fit, Model Fit Results, BMD Estimates, Low Dose Extrapolation, Response at Dose, and Share & Compare. The 'Data Input' step is currently selected and highlighted in blue. To the right of the progress bar are two buttons: 'Preview' and 'DREAM Assist'. Below the progress bar, the section 'Select Dataset Type' is visible. It includes a dropdown menu set to 'Continuous (summary)'. Below this is a table titled 'Dataset [Dose N Mean Stdev]' with columns for Dose, N, Mean, and Stdev. The table has eight rows, with the first row containing the values 1, 1, 1, and 1 respectively. Below the table are two buttons: 'Load examples' and 'Example Data Formats'. Further down is a section titled 'Data Details for Informative Priors' with a 'Show Data Details' button. At the bottom left is a green 'Save' button, and at the bottom right is a blue 'Next' button.

Figure 2.1. The Start Page of a New Analysis

As a note, if you would like to preview the results or view the results from the perspective of someone the analysis was shared with, click the blue button labeled “Preview” in the upper right corner (Figure 2.1). This will take you out of the updating mode. To resume updating the analysis, click the blue button labeled “Actions” in the upper right corner, and then in the drop-down list that appears, click “Update” (Figure 2.2). You can select “Edit a Copy” to create an editable copy of the analysis.

The screenshot shows the Bayesian BMD web application in preview mode. The top navigation bar includes the 'BBMD Bayesian BMD' logo on the left and links for 'Dashboard', 'Help', 'About', 'FAQ', 'Contact', and a user profile icon on the right. Below the navigation bar, the analysis title is 'New Continuous Run Mar 29 2025, 08:37 PM'. The progress bar is identical to Figure 2.1, but the 'Data Input' step is now greyed out, and the 'Response at Dose' step is highlighted in blue. Below the progress bar, the text 'Invalid dataset' is displayed. On the right side, there is an 'Actions' dropdown menu with three options: 'Update', 'Delete', and 'Edit a Copy'.

Figure 2.2. Preview mode

B. Data Input and Pre-analysis

The first tab, which is presented after confirming the selected data type, is the dataset input tab. Select the specific type of data from the dropdown menu, and then put the data into the text box. For continuous and dichotomous datatypes there will be two options in the dropdown menu: Either summary or individual data. If you are doing an analysis with categorical data, there is only one option. Each type of data requires different columns. The following section will explain the types of columns needed for each data type.

Dichotomous summary

If you choose “Dichotomous summary” for dichotomous data, three columns are required for input (from left to right): dose level, total number of subjects and number of subjects affected. The values can be pasted or manually typed, using spaces between values. Different dose groups should be entered in different rows. Below is an example dichotomous summary dataset:

Dose	N (Number of Subjects)	Incidence (Number of Affected)
0	50	1
15.5	49	4
30	50	8
50.6	48	21

Dichotomous individual

If you choose “Dichotomous individual”, two columns are required. The two columns are (from left to right): dose and incidence (either “0” representing no effect or “1” representing with effect). Each row is used for each individual subject. An example of dichotomous individual data is shown below:

Dose	incidence (0 or 1)
0	0
0	0
0	0
0	1
0	0
10	0
10	1
10	0
10	1
10	0
25	1
25	0
25	1
25	1
25	0
50	1
50	1
50	1
50	1

Continuous summary

For “Continuous summary” data type, four columns are needed to describe each dose group: dose, number of subjects, the mean value of response, and the standard deviation of the response. The dataset below is an example of the continuous summary data:

Dose	N (# of Subjects)	Mean	Stdev
0	10	2.82	0.17
100	10	2.91	0.16
200	10	2.95	0.2
400	9	3.22	0.25

Continuous individual

The “Continuous individual” data type only requires two columns (from left to right): dose level and response. The table below shows you an example of this data type.

Dose	Response
0	351.3
0	350.3
0	359.8
0	360.7
0	357.4
2.5	349.8
2.5	352.1
2.5	346.3
2.5	344.7
2.5	350.1
5	340.2
5	341.1
5	345.5
5	331.9
5	347.4
20	331.1
20	320.9
20	319.4
20	308.9
20	314.3

Categorical

The “Categorical” data type requires three columns (from left to right): dose, severity level response. The table below shows you an example of this data type

Dose	Severity	Response
0.00005	0	22
0.00005	1	0
0.00005	2	0
0.00005	3	0
0.01	0	8
0.01	1	0
0.01	2	0
0.01	3	0

Once the data have been entered, press “Save” to save the dataset. Please refresh your browser to make sure the data set has been successfully saved. Once the data set is successfully saved, the data will be visually displayed and summarized in a table as shown in the Figure 2.3. If you hover your mouse over marks on the dose-response plot, you can see more detailed information regarding that point on the plot.



Page 14 of 85

negative) are used to form the p-value. The p-value is equal to the number of contradicting samples divided by the total number of samples. the continuous dataset is appropriate for BMD modeling. Additionally, the NOAEL and LOAEL values will also be displayed. If you uploaded dichotomous data, a trend test (using the Cochran-Armitage trend test, same as the BMDS) is performed. A p-value and z-score will be reported below the data table. The trend test results can be used to judge if the dichotomous dataset is appropriate for BMD modeling. The NOAEL and LOAEL values will also be displayed. If you uploaded categorical data, the NOAEL and LOAEL values for each severity (i.e., category) will be displayed.

To continue to the next tab, press “Next” in the lower right corner.

C. MCMC Settings

On this tab (shown in Figure 2.4), you can specify some settings for the MCMC algorithms.

Bayesian BMD

Dashboard Help About FAQ Contact ?

Manual Example

Progress bar: Data Input, MCMC Settings, Model Settings, Prior Settings, Execute Model Fit, Model Fit Results, BMD Estimates, Low Dose Extrapolation, Response at Dose, Share & Compare

MCMC Settings

Markov Chain Iterations (per chain)

Between 10,000-50,000, inclusive.

Warmup Percent (%)

Between 10-90%, inclusive. This percent of iterations are discarded from the beginning of each chain, and not used for estimating the model distributions.

Number of Markov Chains

Between 1-3, inclusive.

Random Seed

Between 0-99,999, inclusive.

Back Next

Figure 2.4. MCMC settings

i. How to make and change settings

There are four different values that need to be specified in this tab. First, specify the number of Markov chain iterations, between 10,000 and 50,000 (inclusive) iterations per chain. Enter your value into the “Markov Chain Iterations” text box. Next, you need to specify the warmup percentage for each Markov Chain. This is the percentage of iterations discarded from the beginning of each chain; Therefore, those iterations will not be used for estimating model distributions. Put this percentage in the “Warmup Percent (%)” text box. Third, specify the number of Markov chains used in the analysis. Enter a number 1 to 3 (inclusive) into the “Number of Markov Chains” text box. Each chain will use the number of iterations previously specified. The final value is the random seed which is used for reproducing analysis results. The random seed can be 0 to 99,999 (inclusive). Enter this value in the “Random Seed text box”.

Once these values are specified, click “Next” to save the MCMC settings and move to “Model Settings”. Default settings are generally acceptable. However, results in the next step will provide important information that can help you judge if the MCMC settings are appropriate. Based on our testing, **the default settings are adequate for most of the commonly seen dose-response shapes, so we suggest you use the default settings for your initial run.**

ii. How MCMC settings may impact the results

“**Iterations**” is the length of MCMC chain, i.e., the number of posterior samples in each MCMC chain. Default value is 30,000. The allowable range is any integer between 10,000 and 50,000.

“**Number of chains**” is the number of Markov Chains to be sampled. Default value is 1. Allowable range is 1 - 3.

“**Warmup percent (%)**”, the percent of sample in each Markov Chain will be discarded from the final posterior sample. Default value is 50% with an allowable range of 10% - 90%.

So, using the default values, the final number of posterior sample (without the warmup sample) you can get is:

$$30000 \times 1 \times (1 - 50\%) = 15000$$

“**Seed**” is random seed number used in the MCMC algorithms. The number is randomly generated, but you can specify the number for the purpose of reproduction.

D. Model settings

Once data input and MCMC settings are completed, click “Next” to go to the next tab, “Model settings”. In this step, you should choose the model(s) to fit the data. To add a dose-response model, click the plus icon and then click “Create Model” or “Standard Models” on the left panel, shown in Figure 2.5. “Standard Models” includes all default models for each type of data which are described in section i-iii.

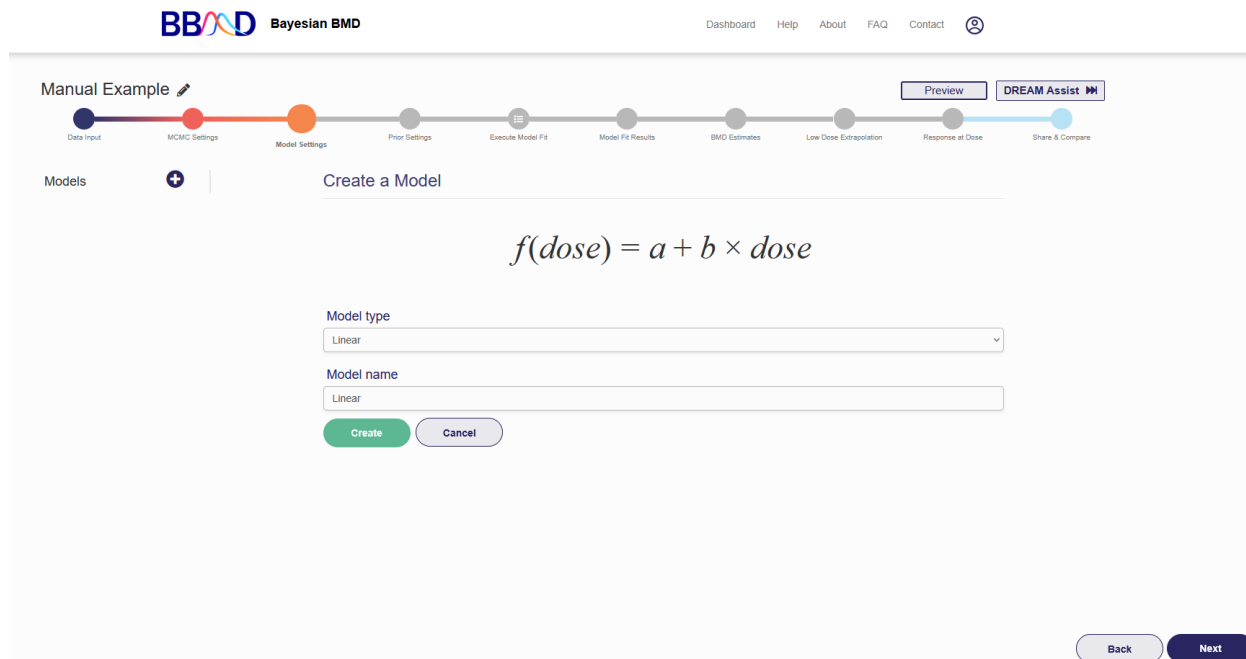


Figure 2.5. “Model Settings” tab

Follow these four steps to add a dose-response model to the list of models for the analysis:

- 1) Click the plus button in the left panel
- 2) Click “Create Model” in the box that appears
- 3) Select one model from the pull-down menu, then give an identifiable name to the model. For models with a power parameter, you need to choose a restriction value for the power parameter. There are five options available in the current system: 0, 0.25, 0.5, 0.75, and 1. The default value is 1.
- 4) Press the “Create” button to add the model to the list of models on the left panel.

To add another model, repeat the steps 1) to 4). The same model with different settings (e.g., the restriction value put on the power parameter) can be added again as a separate model.

To update a model from the model list, follow these 3 steps:

- 1) Click the three dots next name of the model you want to update or delete in the list of models on the left panel.
- 2) If you want to change the current settings of the model, click “Edit Model”, then you can modify the model settings like model type, model name, and the restriction value
- 3) Click “Update” below the settings to update the setting

You can delete a model which is already in the list of models by clicking the three dots next to the model’s name, and then clicking “Delete Model” and “Delete” in the window that pops up (Figure 2.6).

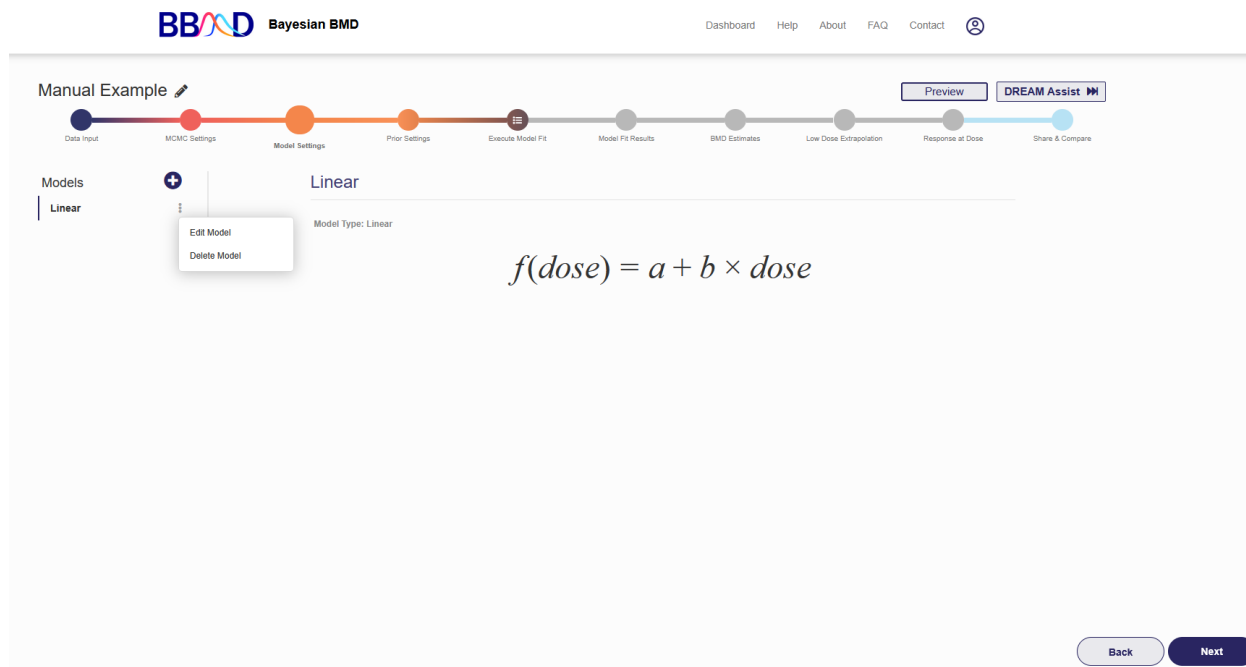


Figure 2.6. Update or Delete a Model in the List

Additionally, for both continuous and dichotomous data, you can choose to use empirical informative priors for the model parameters derived from the general database in Prior Settings. The default setting is the option “Non-Informative Prior” for the model(s) you have selected, which is followed by two options: “Empirical Informative Prior” and “Data Driven Informative Prior”. To turn informative priors on, select “Empirical Informative Prior”. The model(s) you choose to fit the data in Model Settings will be listed in the left panel. When specifying a model in the left panel, empirical prior distribution for each parameter in that model will be displayed. You can edit the parameters of the prior distribution such as mu and sigma for a normal distribution, or beta and alpha for the beta, and gamma distributions based on the recommended ranges shown below after clicking the three dots icon next to the model name and selecting “Edit Model”. Once the distributions are set, distribution curves for the parameters will also be displayed with the dose-response model formula, as shown in Figure 2.7. Currently, the option “Data Driven Informative Prior” is inactivated.

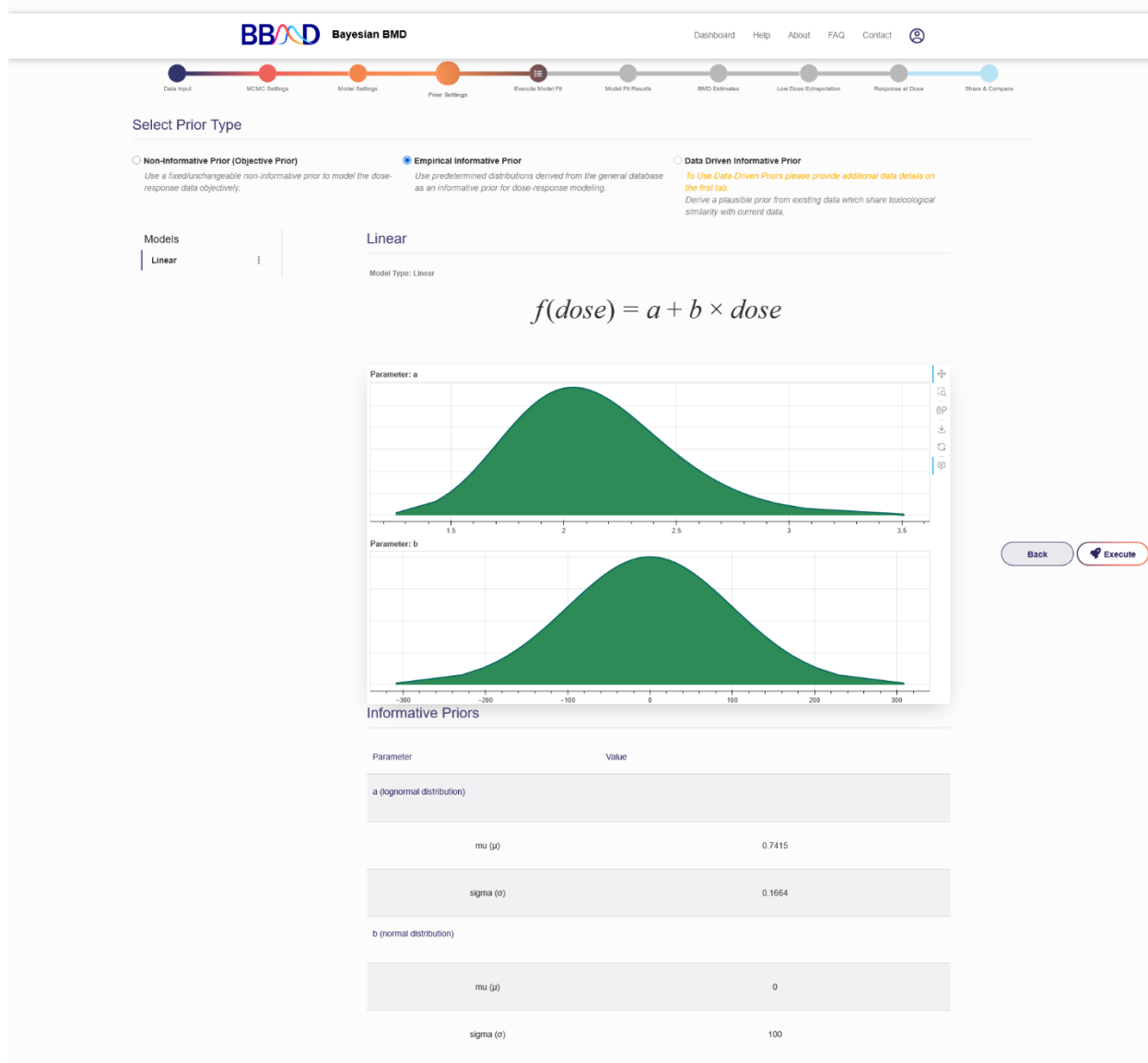


Figure 2.7. Distribution curves for dose-response model's informative priors

When you keep the default setting, non-informative prior distribution is used as the default distribution for all the dose-response model parameters. The non-informative prior distributions for each model are listed below:

i. Dose-Response Models for Dichotomous Data

For Dichotomous data, there are eight models.

1) Quantal Linear Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose})$$

$$a \sim Uniform(0, 1); \quad b \sim Uniform(0, 100)$$

2) Probit Model:

$$f(dose) = \Phi(a + b \times dose)$$

$$a \sim Uniform(-50, 50); b \sim Uniform(0, 100)$$

3) Logistic Model:

$$f(dose) = \frac{1}{1 + e^{(-a - b \times dose)}}$$

$$a \sim Uniform(-50, 50); b \sim Uniform(0, 100)$$

4) Weibull Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-c \times dose^b})$$

$$a \sim Uniform(0, 1); b \sim Uniform(0, 50); c \sim Uniform(restriction, 15)$$

Where 'restriction' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

5) Multistage (2nd Order) Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose - c \times dose^2})$$

$$a \sim Uniform(0, 1); b \sim Uniform(0, 100); c \sim Uniform(0, 100)$$

6) LogLogistic Model:

$$f(dose) = a + \frac{1 - a}{1 + e^{-c - b \times \log(dose)}}$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(-5, 15)$$

7) LogProbit Model:

$$f(dose) = a + (1 - a) \times \Phi(c + b \times \log(dose))$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(-5, 15)$$

8) Dichotomous Hill Model:

$$f(dose) = a \times g + \frac{a - a \times g}{1 + e^{-c - b \times \log(dose)}}$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(-5, 15); g \sim Uniform(0, 1)$$

These eight models are also the models which are part of the "Standard Models" for Dichotomous data.

ii. Dose-Response Models for Continuous Data

Background parameter ‘a’

The background parameter “a” in all eight models has the same uniform distribution used as prior which is derived as follow:

The lower bound of the uniform distribution is always 0, and the upper bound is calculated differently for individual data and summary data.

For individual data,

$$a_{upper} = \max(response) \times 2$$

i.e., doubling the largest response value in the input dataset.

For summary data,

$$a_{upper} = (\max(resp.mean) + 2 \times resp.sd_{mean.max}) \times 2$$

Where $\max(resp.mean)$ is the maximum mean response across all dose groups in the input dataset, and $resp.sd_{mean.max}$ is the response standard deviation in that dose group with the maximum mean response.

Slope parameter ‘b’

For the Linear, Power, Michaelis Menten and Hill model, the lower or upper bound of the parameter b (a slope-equivalent parameter) are determined by the dose-response trend and the overall slope in the input data.

For individual input data and increasing trend:

$$b_{lower} = 0$$

$$b_{upper} = \frac{Max(resp) - Min(resp)}{Dose_{Max_resp} - Dose_{Min_resp}} \times 5$$

For individual input data and decreasing trend:

$$b_{lower} = \frac{Min(resp) - Max(resp)}{Dose_{Min_resp} - Dose_{Max_resp}} \times 5$$

$$b_{upper} = 0$$

Where $Max(resp)$ and $Min(resp)$ are the maximum and minimum response value in the input dataset. And $Dose_{Max_resp}$ and $Dose_{Min_resp}$ are the dose levels corresponding to the maximum and minimum responses respectively.

For summary input data:

$$b_{slope} = \frac{Mean_{Max_dose} + 2 \times SD_{Max_dose} - Mean_{Min_dose} - 2 \times SD_{Min_dose}}{Max(dose) - Min(dose)}$$

Where $Mean_{Max_dose}$ and SD_{Max_dose} are the mean and standard deviation of responses at the maximum dose level, and $Mean_{Min_dose}$ and SD_{Min_dose} are the mean and standard deviation of responses at the minimum dose level. $Max(dose)$ and $Min(dose)$ are the maximum and minimum dose levels in the input dataset. Because dose levels are first normalized to the scale between 0 and 1, $Max(dose)$ is very likely 1 and $Min(dose)$ is very likely 0. Then the prior distribution of the parameter “s” is $Uniform(0, 5 \times b_{slope})$ for increasing trend and $Uniform(5 \times b_{slope}, 0)$ for decreasing trend.

For continuous data, there are eight models.

1) Linear Model:

$$f(dose) = a + b \times dose$$

$$a \sim Uniform(0, a_{upper}); b \sim Uniform(b_{lower}, b_{upper})$$

2) Power Model:

$$f(dose) = a + b \times dose^g$$

$$a \sim Uniform(0, a_{upper}); b \sim Uniform(b_{lower}, b_{upper}); g \sim Uniform(restriction, 15)$$

Where ‘**restriction**’ is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

3) Michaelis-Menten Model:

$$f(dose) = a + \frac{b \times dose}{c + dose}$$

$$a \sim Uniform(0, a_{upper}); b \sim Uniform(b_{lower}, b_{upper}); c \sim Uniform(0, 15)$$

4) Hill Model:

$$f(dose) = a + \frac{b \times dose^g}{c^g + dose^g}$$

$$a \sim Uniform(0, a_{upper}); b \sim Uniform(b_{lower}, b_{upper}); c \sim Uniform(0, 15); g \sim Uniform(restriction, 15);$$

where ‘**restriction**’ is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

5) Exponential 2 Model:

$$f(dose) = a \times e^{b \times dose}$$

$$a \sim Uniform(0, a_{upper}); b \sim Uniform(0, 50) \text{ for increasing trend or } Uniform(-50, 0) \text{ for decreasing trend.}$$

6) Exponential 3 Model:

$$f(dose) = a \times e^{b \times dose^g}$$

$a \sim Uniform(0, a_{upper})$;

$b \sim Uniform(0, 50)$ for increasing trend or $Uniform(-50, 0)$ for decreasing trend;

$g \sim Uniform(restriction, 15)$;

where '**restriction**' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

7) Exponential 4 Model:

$$f(dose) = a \times (c - (c - 1) \times e^{-b \times dose})$$

$a \sim Uniform(0, a_{upper})$; $b \sim Uniform(0, 100)$; $c \sim Uniform(0, 1)$ for decreasing trend or

$c \sim Uniform(1, 15)$ for increasing trend.

8) Exponential 5 Model:

$$f(dose) = a \times (c - (c - 1) \times e^{-(b \times dose)^g})$$

$a \sim Uniform(0, a_{upper})$; $b \sim Uniform(0, 100)$; $c \sim Uniform(0, 15)$ for decreasing trend or

$c \sim Uniform(1, 15)$ for increasing trend; $g \sim Uniform(restriction, 15)$;

where '**restriction**' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

For continuous data, all eight of these models are included in the "Standard Models".

iii. Dose-Response Models for Categorical Data

1) Logistic Model:

$$f(dose) = \frac{1}{1 + e^{(-a - b \times dose)}}$$

$a \sim Uniform(-50, 50)$; $b \sim Uniform(0, 100)$

2) Probit Model:

$$f(dose) = \Phi(a + b \times dose), b \geq 0$$

$a \sim Uniform(-50, 50)$; $b \sim Uniform(0, 50)$

3) Cloglog Model:

$$f(dose) = 1 - e^{-e^{a + b \times dose}}$$

$a \sim Uniform(-100, 100)$; $b \sim Uniform(0, 100)$

4) Quantal Linear Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose})$$

$$a \sim Uniform(0, 1); b \sim Uniform(0, 100)$$

5) Multistage (2nd Order) Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose - c \times dose^2})$$

$$a \sim Uniform(0, 1); b \sim Uniform(0, 100); c \sim Uniform(0, 100)$$

6) Weibull Model:

$$f(dose) = a + (1 - a) \times (1 - e^{-c \times dose^b})$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(0, 50);$$

where '**restriction**' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

7) LogLogistic Model:

$$f(dose) = a + \frac{1 - a}{1 + e^{-c - b \times \log(dose)}}$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(-5, 15);$$

where '**restriction**' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

8) LogProbit Model:

$$f(dose) = a + (1 - a) \times \Phi(c + b \times \log(dose))$$

$$a \sim Uniform(0, 1); b \sim Uniform(restriction, 15); c \sim Uniform(-5, 15);$$

where '**restriction**' is a user defined value and can be 0, 0.25, 0.5, 0.75 or 1.

These eight models are available to be used as dose-response models for categorical data, but the "Standard Models" are reduced to the Logistic, Probit, and Cloglog models.

The models shown on the left panel are the models will be analyzed by the system. Once you are happy with the models selected and the model settings, click "Execute" in the bottom right corner to execute model fitting.

E. Model Fit Results

On the "Model Fit Results" tab, the model fitting results obtained from the previous step are displayed. Click the name of one of the models on the left panel, then the results will be shown on the right (as shown in Figure 2.8) These results include the textual output of model parameter estimation, dynamic dose-response plot, posterior predictive p-value, model weight, correlation matrix, and graphical output of posterior sample of the model parameters (hidden by default). When click "Hide Parameters", the parameter charts for each parameter in the model are displayed as shown in Figure 2.9.

Manual Example [Preview](#) [DREAM Assist](#) 

Model Fit Summary

Seed used: 53902

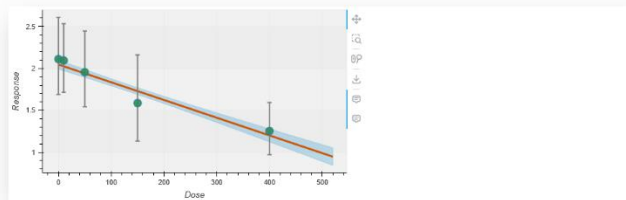
Linear

Linear fit summary

Model Formula

$$f(dose) = a + b \times dose$$

	Mean	MCSE	Stddev	2.5%	5%	50%	95%	97.5%	ILoff	L_het
a	-2.046400	0.000360	0.033335	1.901960	1.902130	-2.046200	-2.180560	-2.112560	0.241	0.999999
b	-0.045219	0.000960	0.002123	-0.095871	-0.046300	-0.044850	-0.742041	-0.722312	0.619	1.000170
lp_	135.696000	0.016220	1.203168	132.306000	133.174000	136.024000	137.066000	137.130000	6207	0.999990
sigma	0.130101	0.000182	0.010072	0.112920	0.115500	0.130200	0.140434	0.152734	9790	0.999999

Assumed distribution of response **Log Normal**Posterior predictive p-value for model fit: **0.415**Model weight: **100%**[Back](#)[Next](#)

Correlation Matrix

	a	b	sigma
a	1	-0.67	0.006620
b	-0.67	1	0.01030
sigma	0.006620	0.01030	1

Parameter Charts

[Show Parameters](#)

Figure 2.8. Results Shown on the “Model fit Results” Page

Parameter Charts

 Hide Parameters

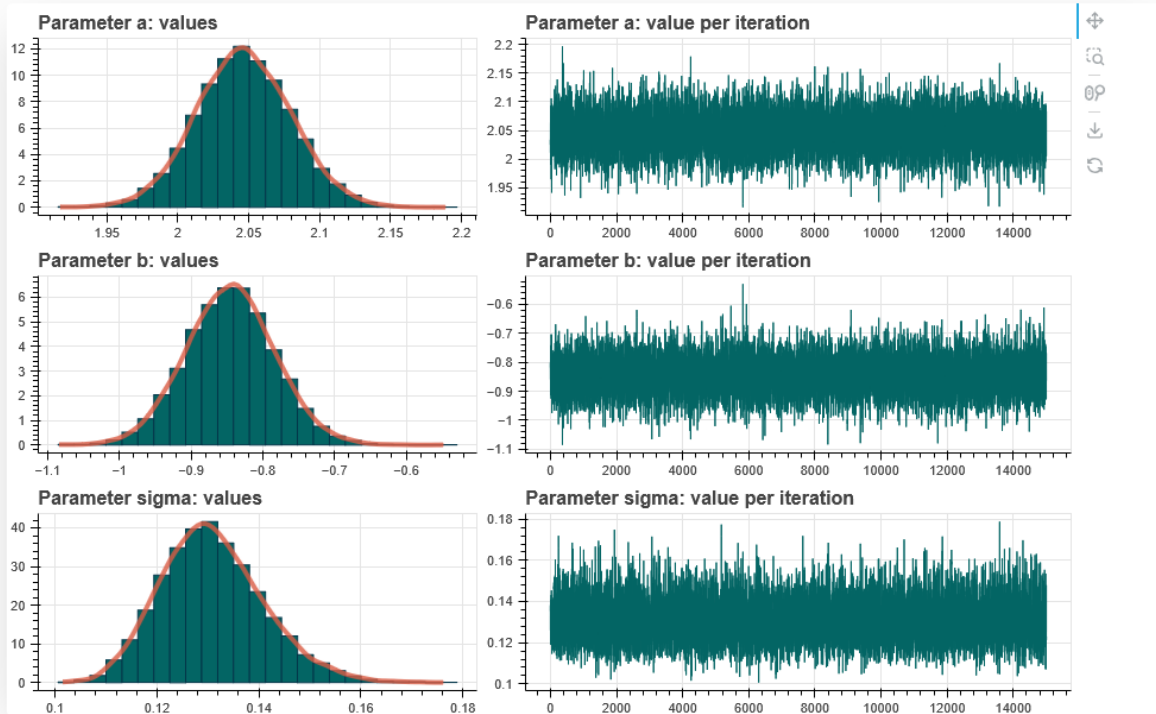


Figure 2.9. Parameter Charts

i. Parameter estimation results

The parameter estimation results, displayed in a table under the model formula, show the statistical summary for the estimated posterior distributions of parameters in the given dose-response model. These results are obtained directly from PyStan's fit output, including some important statistics for model parameters and diagnostic indicators for the MCMC algorithms. The mean, standard error of the mean (MCSE), standard deviation (StdDev), various quantiles (2.5%, 25%, 50%, 75%, and 97.5%), and quantities indicating effective sample size (N_Eff) and chain convergence (Rhat) for each model parameter derived from the posterior distribution of each parameter, as well as information regarding the MCMC execution are summarized in the table. As a note, the "Rhat" can be used to judge if the MCMC chains have converged properly. If the Rhat value is larger than 1.05, you may consider increasing the length of MCMC chains to get better convergence¹.

ii. Posterior Predictive P-Value

¹ Detailed explanation on the Stan outputs can be found at: <https://github.com/stan-dev/stan/releases/download/v2.9.0/stan-reference-2.9.0.pdf>

A posterior predictive p-value (PPP value) is reported below the dynamic dose-response plot. The PPP can be approximated by counting the predicted responses that satisfy the inequality out of the entire posterior sample space. This indicator can be used to judge if the fitting of this particular model is adequate. A large or small p-value means that a discrepancy in predicted data is very likely, further indicating a poor fit. Practically, if the PPP value is between 0.05 and 0.95, then the fitting is adequate. The calculation procedure of PPP value is briefly described below:

- (1) Use each bundle of parameters in the kept posterior sample to form a dose-response model and randomly generate case numbers, y^{rep} , at all dose levels in the original dataset
- (2) Use posterior sample of model parameters to calculate a test statistic for both the original data set (d, n, y) and the replicated data set (d, n, y^{rep}) . The test statistic used in this system is log-likelihood. For parameter values from l -th iteration, we have statistic $T(y, \theta^l)$ and $T(y^{rep}, \theta^l)$.
- (3) For $l = 1, \dots, L$ (the length of posterior sample), compare each pair of $T(y, \theta^l)$ and $T(y^{rep}, \theta^l)$, and count the number of $T(y, \theta^l) > T(y^{rep}, \theta^l)$, say M
- (4) The posterior predictive P-value is $\frac{M}{L}$

A detailed explanation on this procedure can be found in the Chapter of “Model checking and improvement” in *Bayesian Data Analysis* (Gelman et al).

iii. Posterior Model Weight

A model weight ($\hat{m}_j$) for model j is calculated for each model included in the analysis as a statistic for cross-model comparison. The model weight was introduced by (Wasserman, 2000), using the following two equations. The $\hat{m}_j$ value of each selected model j is calculated as follows:

$$\hat{m}_j = \exp \left(\hat{\ell}_j - \frac{d_j}{2} \log(n) \right),$$

where $\hat{\ell}_j$ is a loglikelihood value estimated using one set of posterior samples of model parameters of the j th model, d_j is number of parameters in the j th model, and n is the sample size in the data set.

When all models in the analysis have an equal prior weight, the posterior model weight of model j is calculated by m value estimated from model j divided by the sum of m values estimated from all models in the analysis as the following equation.

$$\Pr(\mathcal{M}_j | Data) = \frac{\hat{m}_j}{\sum_{t=1}^T \hat{m}_t}$$

This function assumes equal model priors for all models selected, so the weight mainly indicates how well the model fits the data. To make the weight more reliable, we use 1000 sets of randomly selected posterior samples of model parameters to calculate the model weights. This

model weights are further applied to the model averaged BMD calculation in the F. BMD Estimation section.

iv. Interactive Dose-Response Plot

A dynamic dose-response plot is shown below the text box. This plot includes original dose-response data and a fitted curve with its 90th percentile interval shaded in blue. When you move your mouse over the dose-response curve, the estimated median and the 5th and 95th percentiles at a particular dose level will display. When you move your mouse over a data point from your inputted dataset, the dose, N, incidence, and the response percentile will also be displayed. Other information displayed in this figure includes the PyStan version, the lower bound placed on the power parameter (if applicable), the posterior predictive p-value (PPP value) for model fit and model weight for cross-model comparison.

v. Correlation Matrix

The fourth item displayed is the correlation matrix for the different model parameters. The correlation matrix is to show the correlation coefficients between different model parameters and is calculated using posterior samples.

vi. Plots for parameter posterior sample

If you click the “Show Parameters” under Parameter Charts, two plots (posterior sample trace plot and estimated probability density plot) will be displayed for each of the parameters in this dose-response model.

Basically, this is the results display tab, meaning that you can only review the results, not give the system additional inputs to modify the results.

F. BMD Estimation

On this page, you can calculate the BMD estimates of your interest. The settings for BMD calculation are slightly different between the analysis for dichotomous data, continuous data and categorical data, therefore, they will be introduced separately.

i. Dichotomous Data

a. How to execute a BMD Analysis

Figure 2.10 is a screenshot of the “BMD estimates” tab for a dichotomous dataset. To create a BMD analysis, you need to follow the four steps below:

- (1) Name the BMD analysis using an easily identifiable name in the “BMD setting name” box.
- (2) Specify a BMR value in the “Benchmark response value” box.
- (3) Give prior model weight to the models included in this analysis. For your reference, “Additional Info” section provides you further info on the performances of each model fitting. The prior weight, PPP value and the calculated weights for each model are provided when you change the slider in the right side of the page below “Additional Info” to

- green. The prior weight will influence the estimation of model averaged BMD value. Giving 0 prior weight to a particular model can exclude the model from model-averaged BMD calculation. The sum of the weights assigned to the individual models are not necessarily required to be 1. The system will automatically convert them. For example, if you give 1 to Logistic model and 3 to Probit model, the system will convert these values to 25% prior weight for the Logistic model and 75% prior weight for the Probit model.
- (4) Click the “Execute” button to execute the BMD analysis using the settings just specified.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

EXAMPLE Data Input MC/MC Settings Model Settings Prior Settings Execute Model Fit Model Fit Results BMD Estimates Low Dose Extrapolation Risk at Dose Share & Compare

BMD + Update BMD Settings

BMD setting name
BMR = 10%

Benchmark response value
0.1
Suggested maximum of the highest incidence rate in the dataset: 0.54; any value in range (1.0000e-4, 0.98) (exclusive) is permitted.

Model-weight priors Additional Info ☒

Model	Prior weight
Logistic	0.125
LogLogistic	0.125
Probit	0.125
LogProbit	0.125
Quantal linear	0.125
Multistage (2nd order)	0.125
Weibull	0.125
Dichotomous Hill	0.125

Enter a non-negative value for each model. Weights will be normalized to sum to 1. If a value is set to 0, then it will not be included in the model-average. If the value is set to the default value (1/N, where N=number of models), then weights will be based purely on model-fit.

Execute **Cancel** **Back** **Next**

Figure 2.10. BMD Analysis Input for Dichotomous Data

Once the BMD analysis is successfully created, the name of the analysis will show on the left panel and the results will be displayed on the right panel, as shown in Figure 19. The results for a BMD analysis will include all the BMD distribution plots from each model, and a summary table. The estimation plots are shown in Figure 2.11, and the summary table is shown in Figure 2.12. To add a new BMD analysis, click the plus button in the left column and repeat steps (1) to (4) above. To edit or delete an existing BMD analysis, click the pencil in the upper right corner. From here you can change the name, BMR value, or model weights the same way as when the analysis was created. To save the changes made, click “Execute” on the bottom of the page. To delete this analysis, click “Delete”. To cancel the modification, click “Cancel”.

EXAMPLE



Preview

BMD

BMR = 10%

BMR = 10%

BMD estimates

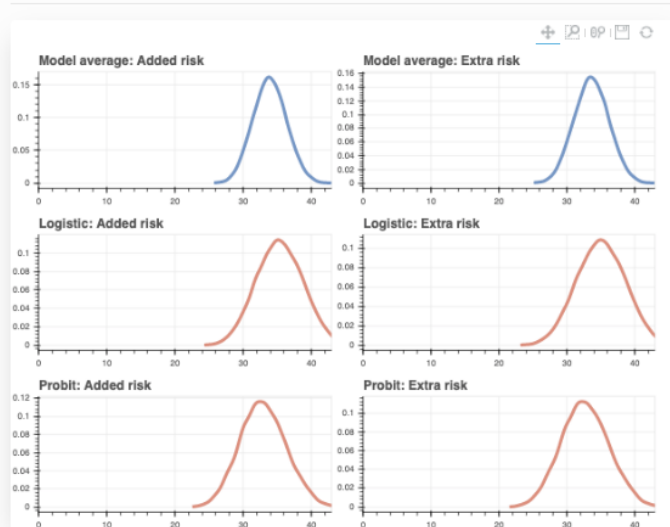


Figure 2.11. BMD Estimation plots from “BMD Estimates” tab

EXAMPLE



Preview

DREAM Assist

BMD

BMR = 10%

BMR = 10%

BMD summary table

Added risk

Statistic	Model average	Logistic	LogLogistic	Probit	LogProbit	Quantal linear	Multistage (2nd order)	Weibull	Dichotomous Hill
Prior model weight	N/A	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250
Posterior model weight	N/A	0.2127	0.1368	0.2418	0.1626	0.0018	0.0546	0.1113	0.0783
Fraction with BMDs	100%	100%	100%	100%	100%	100%	100%	100%	100%
BMD (median)	35.37	35.26	37.48	32.95	37.79	15.28	26.05	37.03	41.02
BMDL (5 th percentile)	26.48	29.79	28.16	27.58	29.10	11.77	18.86	27.21	29.77
25 th percentile	31.87	32.87	33.84	30.62	34.15	13.68	23.50	33.14	36.67
Mean (SD)	35.46 (5.60)	35.34 (3.50)	37.43 (5.57)	32.97 (3.35)	37.73 (5.22)	15.62 (2.77)	25.78 (3.82)	36.95 (5.97)	39.99 (5.27)
75 th percentile	39.06	37.68	41.17	35.21	41.28	17.21	28.33	40.81	44.01
BMDU (95 th percentile)	44.79	41.28	46.39	38.61	46.22	20.60	31.60	46.62	46.88

Figure 2.12. BMD Estimation Results Shown on the “BMD Estimates” tab

b. Explanation of the analysis calculation

For dichotomous data, the BMD will be calculated for both the added risk and extra risk. For the two risks, the BMDs are defined by the following equations, respectively:

$$\text{Added risk: } f(BMD) - f(0) = BMR, \quad \text{Extra risk: } \frac{f(BMD) - f(0)}{1 - f(0)} = BMR;$$

where $f(\cdot)$ represents a dichotomous dose–response model. BMR stands for benchmark response, which is a specified increase in the probability of response and is commonly set at 10%, 5%, or 1%.

In BBMD, the posterior distribution of BMD estimation is established. With the posterior sample, a number of statistics (including the mean, median, standard deviation, and other quantiles) of BMD can be computed and are reported in the “BMD summary table” on the “BMD Estimates” tab. Based on our testing, the median value of the BMD posterior sample is the most reliable estimate for BMD owing to its resistance to some extreme values in the sample. In addition, the 5th percentile of the posterior sample is considered the lower bound of BMD (i.e., BMDL) corresponding to the lower bound of the one-sided 95th confidence interval. The BMDL is usually used as the point of departure for low-dose extrapolation and is therefore of great regulatory interest. The BMD and BMDL values are highlighted in the “BMD summary table”. The same procedures used for determining BMD and BMDL are also applied to continuous and categorical data.

ii. Continuous Data

a. How to execute a BMD Analysis

For continuous data, the BBMD system provides two ways to define BMD: (1) based on central tendency and (2) based on tails (hybrid approach). The basic steps to create a BMD analysis for continuous data are almost identical to the procedure for dichotomous data, except the settings for the BMR. “Central tendency” is the default option for BMD estimation method of continuous data.

Figure 2.13 is a screenshot of the “BMD Estimates” tab for a continuous dataset using Central tendency as the BMD estimation method. To create a BMD analysis, you need to follow the six steps below:

- (1) Name the BMD analysis using an easily identifiable name in the “BMD setting name” box.
- (2) Specify “Central Tendency” as the estimation method to be used in the “BMD estimation method” drop-down menu. If you wish to use “Hybrid Method (tails)”, the detailed steps are in the next section.
- (3) Select an adversity measure from the “Adversity measure” drop-down menu. For BMD defined on central tendency, there are three options for defining the BMR value: a) relative change, b) absolute change, and c) cutoff. Detailed explanations on the three options are in section b.
- (4) Specify an adversity value in the “Adversity value” text box.
- (5) Give prior model weight to the models included in this analysis. “Additional Info” section provides you further info on the performances of each model fitting. The prior weight, PPP value and the calculated weights for each model are provided when you change the slider in

- the right side of the page below “Additional Info” to green. The prior weight will influence the estimation of model averaged BMD value. Giving 0 weight to a particular model can exclude the model from model-averaged BMD calculation. The sum of the weights assigned to the individual models are not necessarily required to be 1. The system will automatically convert them. For example, if you give 1 to Linear model and 3 to Exponential 2 model, the system will convert these values to 25% prior weight for the Linear model and 75% prior weight for the Exponential 2 model.
- (6) Click the “Execute” button to execute the BMD analysis using the settings just specified.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

Manual Example Preview DREAM Assist

BMD + **Update BMD Settings**

BMD setting name
Central tendency: relative change

BMD estimation method
Central tendency

Adversity measure
Relative change
Define adversity via a relative change.

Adversity value
Must be within the range of (0, 1.00).

Model-weight priors Additional Info

Model	Prior weight
Exponential 2	0.125
Exponential 3	0.125
Exponential 4	0.125
Exponential 5	0.125
Hill	0.125
Power	0.125
Michaelis Menten	0.125
Linear	0.125

Enter a non-negative value for each model. Weights will be normalized to sum to 1. If a value is set to 0, then it will not be included in the model-average. If the value is set to the default value (1/N, where N=number of models), then weights will be based purely on model-fit.

Execute Cancel Back Next

Figure 2.13. Settings for Continuous Data BMD Calculation Based on Central Tendency

Figure 2.14 is a screenshot of the “BMD Estimates” tab for a continuous dataset using the hybrid method (tails) as the estimation method. To create a BMD analysis, follow the seven steps below:

- (1) Name the BMD analysis using an easily identifiable name in the “BMD setting name” box.
- (2) Specify “Hybrid Method (tails)” as the estimation method to be used in the “BMD estimation method” drop-down menu.
- (3) Select an adversity measure from the “Adversity measure” drop-down menu. For BMD defined on hybrid method, there are two options for defining the BMR value: a) control group percentile, and b) absolute cutoff value.
- (4) Specify an adversity value in the “Adversity value” text box.

- (5) Specify a BMR value in the “Benchmark response value” text box.
- (6) Give prior model weight to the models included in this analysis. “Additional Info” section provides you further info on the performances of each model fitting. The prior weight, PPP value and the calculated weights for each model are provided when you change the slider in the right side of the page below “Additional Info” to green. The prior weight will influence the estimation of model averaged BMD value. Giving 0 weight to a particular model can exclude the model from model-averaged BMD calculation. The sum of the weights assigned to the individual models are not necessarily required to be 1. The system will automatically convert them. For example, if you give 1 to Linear model and 3 to Exponential 2 model, the system will convert these values to 25% prior weight for the Linear model and 75% prior weight for the Exponential 2 model.
- (7) Click the “Execute” button to execute the BMD analysis using the settings just specified.

BBM Bayesian BMD Dashboard Help About FAQ Contact

Manual Example ✎ Preview DREAM Assist

Data Input MCMC Settings Model Settings Prior Settings Exclude Model Fit Model Fit Results BMD Estimates Low Dose Extrapolation Response at Dose Share & Compare

BMD + **Update BMD Settings**

BMD setting name
Hybrid: Percentile, null

BMD estimation method
Hybrid method (tails)

Adversity measure
Control group percentile
Define adversity via a quantile of the control group.

Adversity value
Must be within the range of (1.0000e-4, 0.50).

Benchmark response value
0.1
Must be within the range of (1.0000e-4, 0.50).

Model-weight priors Additional Info ☐

Model	Prior weight
Exponential 2	0.125
Exponential 3	0.125
Exponential 4	0.125
Exponential 5	0.125
Hill	0.125
Power	0.125
Michaelis Menten	0.125
Linear	0.125

Enter a non-negative value for each model. Weights will be normalized to sum to 1. If a value is set to 0, then it will not be included in the model-average. If the value is set to the default value (1/N, where N=number of models), then weights will be based purely on model-fit.

Execute Cancel Back Next

Figure 2.14. Settings for Continuous Data BMD Calculation Based on Tails

Once the BMD analysis is successfully created, the name of the analysis will show on the left panel and the results will be displayed on the right panel (shown previously in Figure 2.11). To add a new BMD analysis, click the plus button in the left column and repeat steps (1) to (6) above. To edit or delete an existing BMD analysis, click the pencil in the upper right corner. From

here you can change the name, estimation method, adversity measure, adversity value, BMR value, or model weights the same way as when the analysis was created. To save the changes made, click “Execute” on the bottom of the page. To delete this analysis, click “Delete”. To cancel the modification, click ‘Cancel’.

b. Explanation of the analysis calculation

For both types of estimation methods, an adversity measure needs to be specified. For “Central Tendency” there are three options to specify an adversity:

1) Relative change:

For this option, you need to input a value of relative change, e.g., 20%. This means that if the central tendency changes 20% from the control, it will be considered as adverse and the BMD will be calculated accordingly, using the following equation:

$$f(BMD) = f(0) \pm \text{Relative Change} \times f(0)$$

where $f(\cdot)$ represents a continuous dose–response model fit to the central tendency of the data (i.e., the median under the lognormality assumption). The BMD stands for the dose level that satisfies the selected definition equation. The plus/minus sign on the left-hand side is related to the dose-response trend, if increasing, then it is “+”, otherwise it is “-”.

2) Absolute change:

For this option, you need to input a value of absolute change, e.g., 3.2. This means that if the central tendency changes 3.2 from the control, it will be considered as adverse and the BMD will be calculated accordingly, using the following equation:

$$f(0) \pm \text{Absolute Change} = f(BMD)$$

The plus/minus sign on the left-hand side is related to the dose-response trend, if increasing, then it is “+”, otherwise it is “-”.

3) Cutoff:

For this option, you need to input a value of cutoff, e.g., 22.5. This means that if the central tendency is equal to the cutoff value specified, it will be considered as adverse and the BMD will be calculated accordingly, using the following equation:

$$f(BMD) = \text{cutoff}$$

The allowable range for the values of these three options will be automatically calculated based on the trend of the dose-response data and shown.

For the “Hybrid Method (tails)” estimation method, an adversity value must be specified in addition to a BMR value. The hybrid approach considers any response above or below (i.e., corresponding to increasing or decreasing trend) the adversity value as abnormal; thus, the BMD is the dose level where the proportion of the abnormality has increased a certain percent (i.e., BMR) compared with the control. There are two options to specify an adversity:

1) Absolute cutoff value:

For this option, you need to input a value of a cutoff. Then, depending on increasing or decreasing dose-response trend, above or below this value will be considered as adverse.

2) Control group percentile:

For this option, you need to input a percentile value of the control. Then, the below 1st percentile or above 99th percentile of the control distribution is considered as adverse depending on decreasing or increasing.

Mathematically, for increasing trend, the hybrid BMD definition can be expressed as

$$Q(0) - Q(BMD) = BMR \text{ for added risk,}$$

$$\frac{Q(0) - Q(BMD)}{1 - Q(0)} = BMR \text{ for extra risk;}$$

For decreasing trend,

$$Q(BMD) - Q(0) = BMR \text{ for added risk,}$$

$$\frac{Q(BMD) - Q(0)}{Q(0)} = BMR \text{ for extra risk;}$$

where $Q(0)$ represents the quantile of the adversity value at control dose and $Q(BMD)$ represents the quantile of the adversity value at the BMD level.

iii. Categorical Data

a. How to execute a BMD Analysis

Figure 2.15 is a screenshot of the “BMD estimates” tab for a categorical dataset. To create a BMD analysis, you need to follow the five steps below:

- (1) Name the BMD analysis using an easily identifiable name in the “Model name” box.
- (2) Specify a BMR value in the “Benchmark response value” box.
- (3) Select a severity level from the “Severity Level” drop-down list. The number of options equals to levels of the input dataset.
- (4) Give prior model weight to the models included in this analysis. Similarly, “Additional Info” section provides you further info on the performances of each model fitting. The prior weight, PPP value and the calculated weights for each model are provided when you change the slider in the right side of the page below “Additional Info” to green. The prior weight will influence the estimation of model averaged BMD value. Giving 0 weight to a particular model can exclude the model from model-averaged BMD calculation. The sum of the weights assigned to the individual models are not necessarily required to be 1. The system will automatically convert them. For example, if you give 1 to Logistic model and 3 to Cloglog model, the system will convert these values to 25% prior weight for the Logistic model and 75% prior weight for the Cloglog model.
- (5) Click the “Execute” button to execute the BMD analysis using the settings just specified.

Once the BMD analysis is successfully created, the name of the analysis will show on the left panel and the results will be displayed on the right panel (previously shown in Figure 2.11). To add a new BMD analysis, click the plus button in the left column and repeat steps (1) to (5) above. To edit or delete an existing BMD analysis, click the pencil in the upper right corner. From here you can change the name, BMR value, severity level, or model weights the same way as when the analysis was created. To save the changes made, click “Execute” on the bottom of the page. To delete this analysis, click “Delete”. To cancel the modification, click “Cancel”.

The screenshot shows the 'Update BMD Settings' interface for a 'CATEGORICAL EXAMPLE'. At the top, a progress bar indicates the current step is 'BMD Settings'. The main form contains the following sections:

- BMD setting name:** A text input field with the value 'level 3: BMR = 10%'.
- Benchmark response value:** A text input field with the value '0.1'. Below it, a note states: 'Suggested maximum of the highest incidence rate in the dataset: 0.1000, any value in range (1.0000e-5, 0.50) (exclusive) is permitted.'
- Severity level:** A dropdown menu showing 'level 3 (severity #3)'.
- Model-weight priors:** A section with a toggle switch labeled 'Additional Info' (currently turned on). It contains three rows:

Model	Prior weight
Logistic	0.3333333333333333
Probit	0.3333333333333333
Cloglog	0.3333333333333333

At the bottom, there is a note: 'Enter a non-negative value for each model. Weights will be normalized to sum to 1. If a value is set to 0, then it will not be included in the model-average. If the value is set to the default value (1/N, where N=number of models), then weights will be based purely on model-fit.' Below this note are 'Execute' and 'Cancel' buttons. On the right side, there are 'Back' and 'Next' buttons.

Figure 2.15. Settings for a Categorical Data BMD

b. Explanation of the analysis calculation

The BMD is defined based on “Central Tendency” for each severity level. The BMD are calculated for both the added risk and extra risk at each severity level. For the two risks, the BMDs are defined by the following equations, respectively:

$$\text{Added risk: } f(BMD) - f(0) = BMR, \quad \text{Extra risk: } \frac{f(BMD) - f(0)}{1 - f(0)} = BMR;$$

where $f(\cdot)$ represents a dichotomous dose–response model. BMR stands for benchmark response, which is a specified increase in the probability of response and is commonly set at 10%, 5%, or 1%.

iv. Model Averaged BMD Calculation

Before calculating a model averaged BMD, a posterior sample of the BMD from each individual model should be obtained. The process to get the posterior sample was described in the previous three sections. Here, we focus on the method to get model averaged BMD.

a. Posterior Model Weight for Model Averaged BMD Calculation

In this step, the prior model weight specified by users will be used in the posterior model weight calculation. The function is shown below. The $\hat{m}$ for each model is calculated using the same procedure described in the previous section.

$$\Pr(\mathcal{M}_j|Data) = \frac{\hat{m}_j \Pr(\mathcal{M}_j)}{\sum_{t=1}^T \hat{m}_t \Pr(\mathcal{M}_t)}$$

Based on the function above, we know that the posterior weight of a model will be 0 if the prior weight for the model is specified as 0. That is, the model with 0 weight will be excluded from the analysis.

b. Model Averaged BMD Calculation

For each model, we have posterior sample of BMD with the same length as the model parameters. Using default value, we should have:

$$\begin{aligned} &BMD_{1-1}, BMD_{1-2}, \dots, BMD_{1-15000} \text{ for model 1} \\ &BMD_{2-1}, BMD_{2-2}, \dots, BMD_{2-15000} \text{ for model 2} \\ &\dots \end{aligned}$$

Then the posterior sample of model averaged BMD is calculated as a mixture distribution over all models:

$$\Pr(BMD_{ma}|Data) = \sum_{j=1}^J \Pr(BMD_j|M_j, Data) \Pr(M_j|Data),$$

That is:

$$BMD_{MA-15000} = BMD_{1-15000} \times w_1 + BMD_{2-15000} \times w_2 + \dots$$

Therefore, we will have the same size of posterior sample for model averaged BMD. $w_1, w_2, \dots$ are posterior model weight (prior model has been integrated) calculated in the previous section.

G. Probabilistic RfD Estimation

On this page, you can calculate the RfD estimates of your interest using your previously calculated BMD estimates. The RfD estimate settings are the same for each data type.

Figure 2.16 is a screenshot of the “RfD estimates” tab. To create a RfD analysis, you need to follow the eleven steps below:

- 1) Name the RfD analysis in the “RfD analysis name” box
- 2) Specify the dose units in the “Dose Units” box
- 3) (Optional) Choose a random seed for the “Random number seed” box

- 4) Select a BMD analysis to be used as the point of departure. Your BMD analyses from the previous tab will be the possible selections in the drop-down menu.
- 5) Specify the Allometric Scaling settings. Select a test species from the “Test species” drop-down menu. Specify the test species body weight in the “Test species body weight” box. Choose a human body weight in the “Human body weight” box. Lastly, specify the Allometric scaling exponent mean and standard deviation in the two corresponding boxes. For each species you choose the default values for the parameters in “Allometric Scaling” section will automatically be displayed.
- 6) Give an Animal to Human Uncertainty geometric mean and geometric standard deviation. These two values go in their respective boxes.
- 7) Specify Duration of Exposure details to extrapolate non-chronic exposures to chronic exposures. Select the duration of exposure from the drop-down menu. Then give the short-term exposure geometric mean and geometric standard deviation in the corresponding boxes.
- 8) (Optional) Add up to two additional uncertainties. Each additional uncertainty requires a name, geometric mean, and geometric standard deviation. These choices go in the respective boxes.
- 9) Specify the Human Variability geometric mean and geometric standard deviation in the corresponding boxes.
- 10) Give the target population-based incidence in the I* box.
- 11) Click the “Save” button to execute the RfD analysis using the settings just specified.

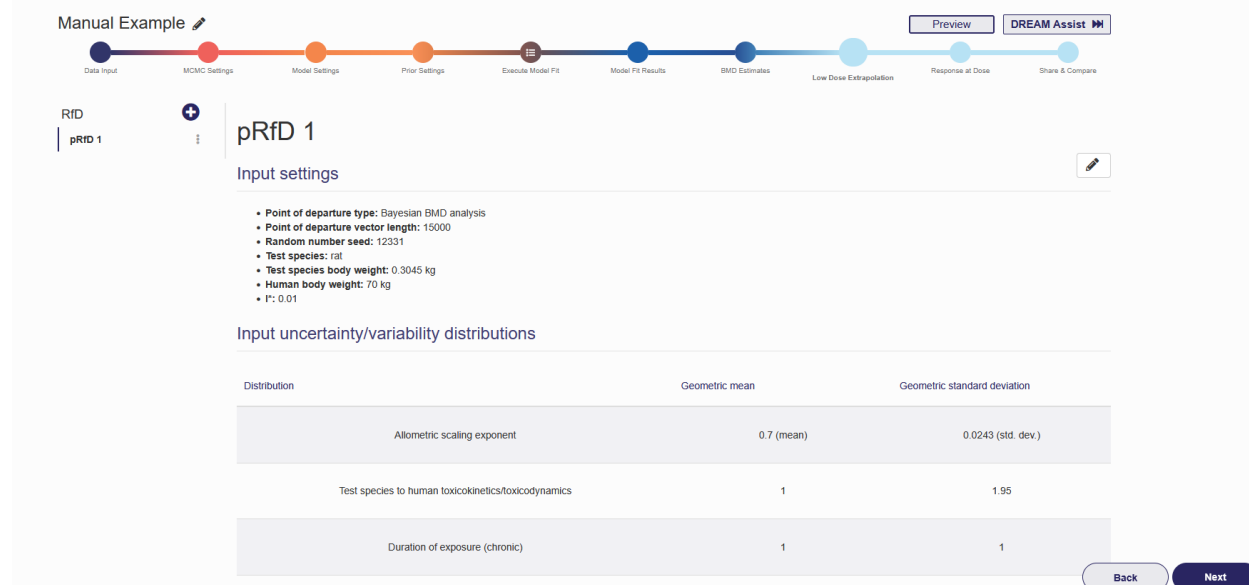


Figure 2.17. RfD Estimate results shown after RfD analysis finishes

H. Risk/Response at Dose

On this page, you can calculate the RAD (Response at Dose) estimates of your interest. The settings for RAD analyses using dichotomous data and continuous data are identical. The only difference from these two data types and categorical data is that a specified severity level is required for categorical data.

Figure 2.18 is a screenshot of the "Risk at Dose" tab (This tab is called "Response at Dose" for continuous and categorical data). To create a new RAD analysis, you need to follow the five steps below:

- 1) Name the RAD analysis using an identifiable name in the "Model Name" box.
- 2) Specify a dose value in the "Dose value" box.
- 3) Only for categorical data – Choose a severity level from the "Severity level" drop-down menu.
- 4) Give prior model weights to the models included in this analysis. Giving 0 weight to a particular model can exclude the model from model-averaged RAD calculation. The sum of the weights assigned to the individual models are not necessarily required to be 1. The system will automatically convert them.
- 5) Click the "Save" button to execute the RAD analysis using the settings just specified.

EXAMPLE

Data input

Model Settings

Prior Settings

Execute Model Fit

Model Fit results

BMD Estimates

Low Dose Extrapolation

Risk at Dose

Share & Compare

Preview

DREAM Assist

Models

Model Name

Dose = 25

Dose Value

25

Dose must be greater than or equal to 0 and less than 112.50

Model-Weight Priors

Model	Prior weight
Logistic	0.125
LogLogistic	0.125
Probit	0.125
LogProbit	0.125
Quantal linear	0.125
Multistage (2nd order)	0.125
Weibull	0.125
Dichotomous Hill	0.125

Enter a non-negative value for each model. Weights will be normalized to sum to 1. If a value is set to 0, then it will not be included in the model-average. If the value is set to the default value (1/N, where N=number of models), then weights will be based purely on model-fit.

Execute

Cancel

Back

Next

Figure 2.18. Risk at Dose analysis for Dichotomous Data

Once the RAD analysis is successfully created, the name of the analysis will be shown on the left panel and the results will be shown on the right as seen in Figure 2.19. To add a new RAD analysis, click the plus button in the left column and repeat steps (1) to (5) above. To edit or delete an existing RAD analysis, click the pencil in the upper right corner. From here you can change all analysis settings the same way as when the analysis was created. To save the changes made, press “Update” on the bottom of the page. To delete this analysis, press “Delete”.

	Dose = 25								
Risk at Dose summary table									
Risk									
Statistic	Model average	Logistic	LogLogistic	Probit	LogProbit	Quantal linear	Multistage (2nd order)	Weibull	Dichotomous Hill
Prior model weight	N/A	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250
Posterior model weight	N/A	0.2127	0.1368	0.2418	0.1626	0.0018	0.0546	0.1113	0.0783
Fraction with Risk	100%	100%	100%	100%	100%	100%	100%	100%	100%
Risk (median)	0.06932	0.06921	0.06132	0.07084	0.05488	0.18	0.12	0.06720	0.04266
5 th percentile	0.05375	0.03648	0.02912	0.03757	0.02393	0.15	0.08565	0.03431	0.01671
25 th percentile	0.06262	0.05407	0.04619	0.05583	0.04002	0.17	0.10	0.05210	0.02986
Mean (SD)	0.06979 (0.01018)	0.07160 (0.02408)	0.06495 (0.02597)	0.07332 (0.02437)	0.05811 (0.02454)	0.18 (0.02446)	0.12 (0.02478)	0.07103 (0.02654)	0.04647 (0.02283)
75 th percentile	0.07636	0.08661	0.07996	0.08857	0.07237	0.20	0.14	0.08598	0.05896
95 th percentile	0.08751	0.11	0.11	0.12	0.10	0.23	0.17	0.12	0.08855

Figure 2.19. RAD Estimate results shown in on the “Risk at Dose” tab

I. Share the Analysis

The final tab for this analysis is the “Share” tab, shown in Figure 2.20. By default, all analyses in a personal account can only be accessed by the owner of the account. If you would like to share your analysis with others, you can change the settings from “Private” to “Send Back”, “Share”, “Send to a Peer”, or “Dream Assist”. The public setting (i.e., “Share”) allows you to send the created URL to others for them to access and review (but not edit) the analysis.

You can also export the results of the analysis into Word or Excel formats. Before exporting results, you can customize the parts of the analysis included on the reports. By clicking “show more” on the right side of the page, the customization options will appear. Here you can specify the BMD results, RAD results, or model summaries to be included. If you wish to change the report settings in the future, you can return to this analysis and export a new report. Exporting the results will send a link to your account’s email where you can download the reports.

At any time during the updating or reviewing stage, if you want to change to another existing analysis, you can click the “Dashboard” button on the top right corner to switch to the summary page for the existing analyses and access another analysis.

The screenshot shows the 'Share & Compare' section of the Bayesian BMD web application. At the top, the navigation bar includes the 'BBMD Bayesian BMD' logo and links for 'Dashboard', 'Help', 'About', 'FAQ', 'Contact', and a user profile icon. The 'Share & Compare' header is followed by a sub-header 'Share & Compare' and a descriptive sentence: 'Use the cards below to share your analysis. You can also compare to EPA's BMDS and download reports below.' Below this are five share options, each with an icon and a description: 'Send Back' (person with checkmark), 'Private' (person icon, highlighted with a blue border), 'Share' (group of people icon), 'Send to a Peer' (group of people icon), and 'DREAM Assist' (DREAM TECH, LLC logo). Under the 'Share' option, there is a 'BBMD Reports' section with three report types: 'BMD Analysis Report (Word Format)', 'Posterior Sample of Model Parameters (Excel Format)', and 'Posterior Sample of BMD Estimates (Excel Format)'. Each report type has a 'Show More' link and a 'Download' button. At the bottom, there is a 'Compare to EPA BMDS' section with a note: 'Please note you must have at least 1 BMR definition and that by clicking this button your data will be sent to the EPA's BMDS online server for executing.' Below the note are two buttons: 'Execute' and 'Back'.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

Share & Compare

Use the cards below to share your analysis. You can also compare to EPA's BMDS and download reports below.

Send Back
Return to the member that requested the DREAM Assist once reviewed.

Private
This analysis can only be viewed by you

Share
This analysis can be viewed by others

Send to a Peer
Send to a reviewer in your organization

DREAM Assist
Share your analysis with DREAM Tech and our team will assist with your analysis

BBMD Reports

- ☐ BMD Analysis Report (Word Format) Show More
- ☐ Posterior Sample of Model Parameters (Excel Format) Show More
- ☐ Posterior Sample of BMD Estimates (Excel Format) Show More

Download

Compare to EPA BMDS

Please note you must have at least 1 BMR definition and that by clicking this button your data will be sent to the EPA's BMDS online server for executing.

Execute Back

Figure 2.20. “Share” tab for a BMD analysis using a single dataset

III. Batch Processing for Analysis

An automatically generated name “New Batch Run *Month Day Year, HH:MM AM/PM*”, is assigned to the newly started analysis. You can click the pencil button next to the analysis name, as seen in Figure 3.1, to make the name more identifiable. Without inputting any data, you will not be able to advance through the tabs. Therefore, the first step in the analysis is to input the batch dose-response data.

The screenshot displays the 'Batch BMD Analysis' workflow. At the top, a progress bar shows steps: Data Input (active), Model Settings, MCMC Settings, BMR Settings, Execute, Results, and Share/Review. Buttons for 'Preview' and 'DREAM Assist' are visible. Below the progress bar, the 'Batch BMD Analysis' section includes a note: 'The analyses provided in this section are designed to provide an efficient but slightly simplified way to analyze a number of datasets all together, including modeling fitting and BMD estimation results.' A 'Dataset type' dropdown menu is set to 'Continuous (Summary)'. Under 'Upload Datasets File', there is a 'Choose File' button and a link to 'Example Dataset'. A 'Description' text area is at the bottom. At the very bottom, there are 'Save' and 'Next' buttons.

Figure 3.1. Start of a new Batch Analysis

A. Data Input

Note: For data input, a “.csv” file needs to be uploaded. You can click “Example Dataset” to see the required format for each type of dataset.

If you choose to do an analysis on dichotomous data, only dichotomous summary data is allowed. Four columns are required for input (from left to right): Dataset Index, Dose level, number of subjects, number of subjects affected.

If you choose to do an analysis on continuous data, both continuous summary and continuous individual data is accepted. For continuous summary data, five columns are required. The columns (left to right) are Dataset Index, dose level, number of subjects, mean, standard deviation. For continuous individual data, three columns are required (from left to right): Dataset Index, dose level, response.

If you choose to do an analysis on categorical data, four columns are required (left to right): Dataset Index, dose level, severity level, response.

After inputting your dataset, you can add a description to the analysis in the “Description” box. When the dataset is successfully uploaded and a description is added, press “Next” in the lower right corner to advance to the next tab.

B. Model Settings

[Dashboard](#)
[Help](#)
[About](#)
[FAQ](#)
[Contact](#)

Need More? With DREAM Premium, you'll get access to more features. Get the first 30 days free.
Get Started

EXAMPLE BATCH

Data Input

Model Settings

MCMC Settings

BMR Settings

Results

Results

Share/Review

Preview
DREAM Assist

Models

Model name	Include	Formula	Power parameter lower bound
Linear	☒	$f(dose) = a + b \times dose$	N/A
Power	☒	$f(dose) = a + b \times dose^c$	1.0
Hill	☒	$f(dose) = a + \frac{b \times dose^c}{c^c + dose^c}$	1.0
Exponential 2	☒	$f(dose) = a \times e^{b \times dose}$	N/A
Exponential 3	☒	$f(dose) = a \times e^{b \times dose^c}$	1.0
Exponential 4	☒	$f(dose) = a \times [c - (c - 1) \times e^{-b \times dose}]$	N/A
Exponential 5	☒	$f(dose) = a \times [c - (c - 1) \times e^{-(b \times dose)^c}]$	1.0

Save
Back
Next

Figure 3.2. “Model Settings” tab

Figure 3.2 shows the “Model Settings” tab, the next step in performing a batch BMD analysis. To include a model in the analysis, ensure that the check box is checked. The model formula is in the next column. For models there is a power parameter lower bound, you can choose a value in the last column. The power parameter lower bound options are 0, 0.25, 0.5, 0.75, and 1.

The models available for each data type are slightly different. For dichotomous data the eight available models are: Logistic, Probit, Quantal Linear, Multistage (2nd Order), Weibull, LogLogistic, LogProbit, Dichotomous Hill. For continuous data the seven available models are: Linear, Power, Hill, Exponential 2, Exponential 3, Exponential 4, Exponential 5. For categorical data the eight available models are Logistic, Probit, Cloglog, Quantal Linear, Multistage (2nd order), Weibull, LogLogistic, LogProbit.

C. MCMC Settings

On this tab you can specify the settings for the MCMC algorithm.

“Iterations” is the length of MCMC chain, i.e., the number of posterior samples in each MCMC chain. Default value is 30,000. The allowable range is any integer between 10,000 and 50,000.

“Number of chains” is the number of Markov Chains to be sampled. Default value is 1. Allowable ranges 1 - 3.

“Warmup percent (%)”, the percent of sample in each Markov Chain will be discarded from the final posterior sample. Default value is 50% with an allowable range of 10% - 90%.

So, using the default values, the final number of posterior sample (without the warmup sample) you can get is:

$$30000 \times 1 \times (1 - 50\%) = 15000$$

“Seed” is random seed number used in the MCMC algorithms. The number is randomly generated, but you can specify the number for the purpose of reproduction.

Once these values are specified, click “Save” to save the MCMC settings.

D. BMR Settings

The BMR settings tab, shown in Figure 3.3, is where you specify the BMR type and BMR values for the analysis. To add a new benchmark dose calculation, click the plus button on the right side of the right panel. When adding a new BMR, specify either Extra or Added for the benchmark dose calculation, then specify the BMR value. To delete a BMR, click the red trash can button on the right side of the table. You can add up to 4 different BMRs for the analysis.

When you are satisfied with the specified BMRs, either press “Save” in the bottom left corner of the page to save the analysis to come back to later, or press “Execute” in the bottom right corner to execute the batch analysis.

Bayesian BMD

Dashboard Help About FAQ Contact

Need More? With DREAM Premium, you'll get access to more features. Get the first 30 days free. Get Started

EXAMPLE BATCH

Progress: Data Input, Model Settings, MCMC Settings, **BMR Settings**, Execute, Results, Share/Review

Benchmark Responses (BMRs)

Assumed Distribution of Response: Normal ☒ Log Normal ☐

BMR Definition	BMR value
None. Click the "+" button above to create a row	

Save Back Execute

Figure 3.3. “BMR Settings” tab

E. Execute

While the BMD Batch analysis is executing, you may leave the system. When the analysis is completed, you will receive an email notification, and a link to your analysis. Then you will be able to view, share, and export the results.

F. Results

This tab, shown in Figure 3.4, displays the model summaries and BMD summaries for the batch analysis. You can either download the Excel spreadsheet versions of the results or view them in the tables on this page.

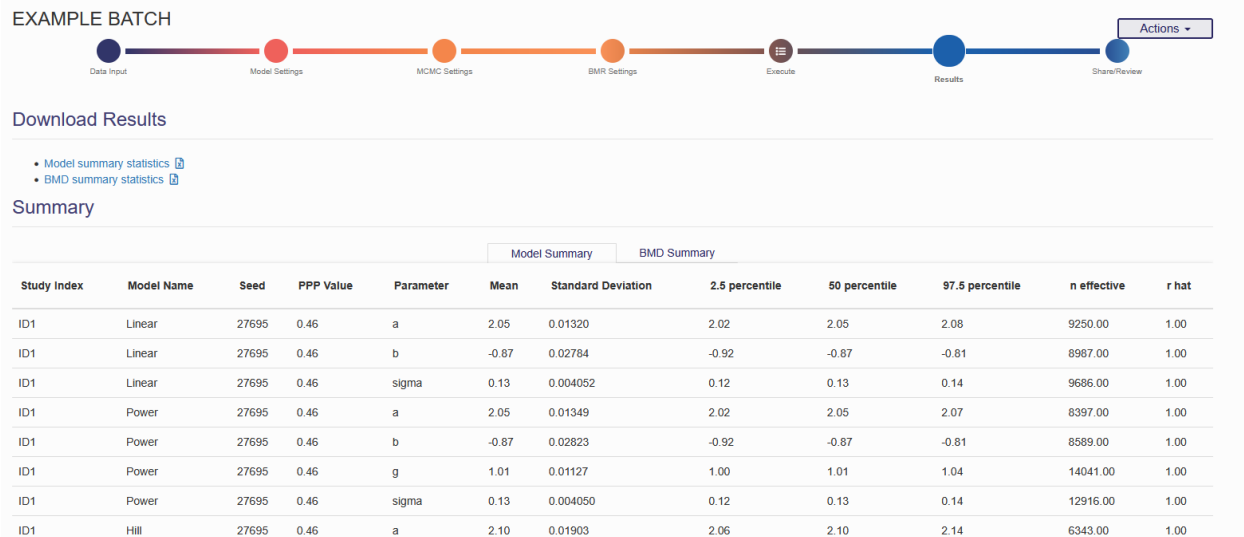


Figure 3.4. “Results” tab for BMD Batch analysis

In the Model Summary table, each model has a row for each parameter in the model. The columns of the table are Study Index (dataset name), Model Name, Seed, PPP Value, Parameter, Mean, Standard Deviation, 2.5 percentile, 50 percentile, 97.5 percentile, n effective, and r hat.

The BMD summary table displays the results for the BMDs specified a couple tabs earlier. The table columns are Study Index (dataset name), BMR Type, BMR Value, Model, model prior weight, model posterior weight, BMDL, BMD, and BMDU.

G. Share/Review

The final tab for this analysis is the “Share” tab, shown in Figure 3.5. By default, all analyses in a personal account can only be accessed by the owner of the account. If you would like to share your analysis with others, you can change the settings from “Private” to “Send Back”, “Share”, or “Dream Assist”. The “Share” setting allows you to send the created URL to others for them to access and review (but not edit) the analysis. The “Send Back” setting allows you to Return to the member that requested the DREAM Assist once reviewed. Lastly, the “DREAM Assist” setting will share this analysis with the DREAM Tech team, who can assist with the analysis. The “DREAM Assist” settings are coming soon.

Like in the previous tab, you can also download the summary Excel spreadsheets for this analysis. There are no differences between the spreadsheets available to download from the “Results” or “Share/Review” tabs.

BBMD

Bayesian BMD

Dashboard

Help

About

FAQ

Contact

Need More?

With DREAM Premium, you'll get access to more features. Get the first 30 days free

Get Started

New Batch Run Mar 30 2025, 01:44 PM

Data Input

Model Settings

MCMC Settings

BMR Settings

Execute

Results

Share/Review

Actions

Analysis Summary

Please review analysis details and publication preferences. By selecting desired sections, you can either include or exclude tabs by selecting the check boxes below

Send Back

Return to the member that requested the DREAM Assist once reviewed.

Private

This analysis can only be viewed by you

Share

This analysis can be viewed by others

COMING SOON!

DREAM Assist

Share your analysis with DREAM Tech and our team will assist with your analysis

Download Results

Model summary statistics

BMD summary statistics

Back

Figure 3.5. “Share” tab for a BMD Batch analysis

Page 47 of 85

IV. BMD Analysis for Genomic Data

A. Introduction to BMD analysis for Genomic Data

This analysis is designed to perform BMD modeling and estimation for genomic data.

When beginning a new analysis, an automatically generated name “New Genomic BMD Analysis *Month Day Year, HH:MM AM/ PM*” is assigned to the analysis. You can click the pencil button next to the analysis name, as seen in Figure 4.1, to make the name more identifiable.

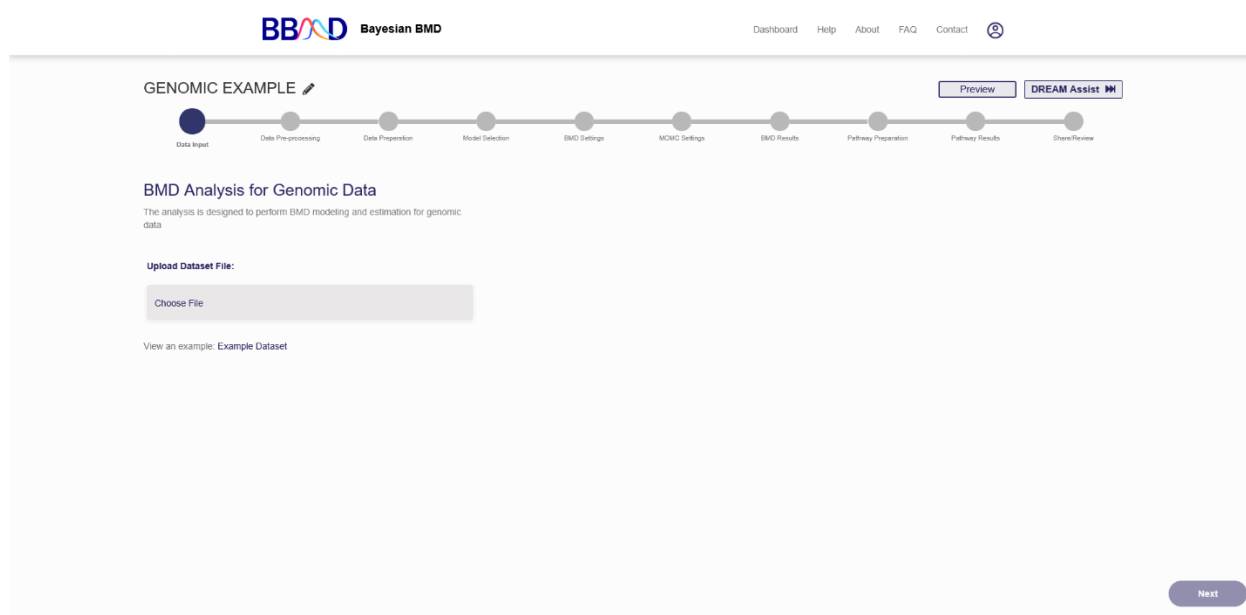


Figure 4.1. First page of a new genomic BMD analysis

B. Data input

i. How to input data into the system

The first step in performing a BMD Analysis for genomic data is to input your genomic dataset. The dataset needs to be inputted as a “.csv” file and contain the proper columns/rows. The first row should be the Sample IDs, followed by the second row which should contain the dose levels. After that, each row should be labeled with the Gene name. Each column should be a different sample ID and dose level. The table below shows the first few columns and rows of an example dataset.

Example Genomic Dataset

SampleID	2D_RG_PLAT_1_09 3016_G23	2D_RG_PLAT_1_09 3016_M23	2D_RG_PLAT_1_09 3016_L21	2D_RG_PLAT_1_09 3016_M21
Dose	0	0	4.742541	4.742541
ACAA1_48	7.907	8.564	7.721	7.873
CYP2C8_15146	10.334	10.81	10.192	10.192

EREG_21427	4.083	4.266	4.898	4.054
IL1B_3325	3.82	3.322	3.783	3.469

⋮

The row labels “SampleID” and “Dose” must be present in the dataset, otherwise the dataset will be considered invalid.

ii. How to interpret the data input results

After a dataset is successfully uploaded, a summary table, principal component plots, and a density plot will be available to view, as shown in Figure 4.2. The summary table will show the dataset file name, the number of doses, the dose levels, the number of samples, number of valid genes, and invalid genes (if any were present). An invalid gene is a gene with insufficient data, such as missing response values. If your dataset contained invalid genes, those genes were removed, and you can continue the analysis with the remaining genes.

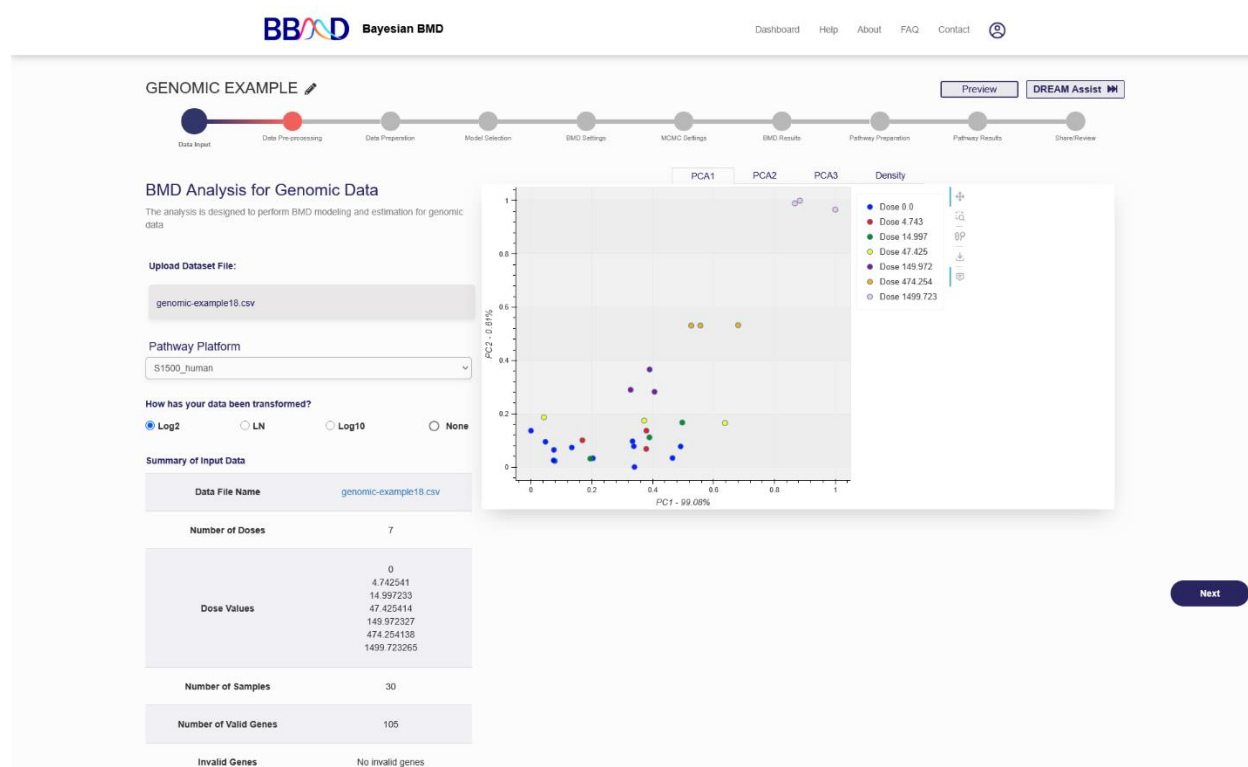


Figure 4.2. Summary results and the plots available after successfully uploading a genomic dataset

There are also four different plots available to display the relationships found within the data. The first three plots are Principal Component plots, showing PC1 vs. PC2, PC1 vs. PC3, and PC2 vs. PC3. For these plots, each dose level is given a unique color, and you can hover over the points on the plot to show their dose level and sample ID. The PCA plots are intended to demonstrate how well the genomic data clustered before performing an analysis. PCA plots can be used to identify the outlier samples that seem “far away” from others in their replicate group. The outliers may reflect some biological interactions that should not be ignored

depending on your experimental system. Combining the PCA plots and the density plot to decide if data should be included in the analysis or not. The last plot is the Density plot, which gives you an idea about the signal distribution across microarray chips. Arrays that have very different distributions in the density plot should be checked carefully. Generally, these arrays will show up as problematic and if so, should be removed before analysis.

Before you can continue onto data pre-processing, you must input the type of transform the data has already undergone. The options are $\log_2$, $\ln$, or $\log_{10}$. If the data has not yet been transformed, select “None”, and a $\log_2$ transform will be applied. The transform type is necessary for the preprocessing algorithms, so it is critical that the correct option is selected.

Once have finished the data transform step and are satisfied with the summary data, you can click “Next” in the bottom right corner to continue onto the “Data Pre-processing” tab.

C. Data Pre-processing

i. How to perform data pre-processing

The next step in performing a genomic BMD analysis is to use the three preprocessing algorithms to rule out genes that don’t fit the specified criteria. This step will reduce the number of genes to be used in the analysis in the BMD analysis and Pathway analysis steps. Figure 4.3 shows the “Data Pre-screening” tab before the preprocessing algorithms have executed.

Bayesian BMD

Dashboard Help About FAQ Contact

GENOMIC EXAMPLE

Progress bar: Data Input, Data Pre-processing, Data Preparation, Model Selection, BMD Settings, MCMC Settings, BMD Results, Pathway Preparation, Pathway Results, Share/Review

Select Preprocessing algorithms

- ☒ One-way Anova
 - Number of permutations: 100
- ☐ Williams Test
 - Number of permutations: 100
- ☐ Origen
 - Number of bootstrap iterations: 500

Specify Preprocessing screening criteria

Specify the Screening P-value: 0.05

Specify Fold Change: 2

Execute

Back Next

Figure 4.3. “Data Pre-processing” tab for a Genomic BMD Analysis

Performing data preprocessing takes place in two steps. The first step is to choose the preprocessing algorithms to use and specify the required settings. The second step is to specify the preprocessing screening settings.

The three algorithms used for preprocessing are “One-way Anova”, “Williams Test”, and “Origen”. All three of the algorithms use the same screening criteria to select genes that pass,

but the screening values are calculated differently in each algorithm, so different genes may pass one algorithm but not another. If you use either “Williams Test” or “Oriogen” you must also specify an additional setting. For the “Williams Test”, you must enter the number of permutations to be used during the algorithm, and for “Oriogen” you must specify the number of bootstrap iterations to be used during the algorithm. Using a higher number of permutations or bootstrap iterations may allow the algorithms to become more precise but with the cost of taking longer to execute.

The screening criteria that you need to specify is the screening P-value and the fold change. For a gene to pass prescreening, the P-value must be below the value entered in the corresponding text box. The gene’s fold change must be above the value entered into that text box.

Once you are satisfied with the selected algorithms, algorithm settings, and prescreening criteria, press “Execute” in the lower right corner. The prescreening algorithms can take a long time to execute, so an email will be sent to your inbox when the prescreening is complete so you can return to the analysis to view the results and move onto BMD and Pathway analyses.

ii. How to interpret the data pre-processing results

When the preprocessing algorithms have finished executing, a summary of the results, and two volcano plots will be displayed, shown in Figure 4.4. The summary of results will show the number of genes that passed each prescreening algorithm. The volcano plots will show the max fold change vs. the adjusted p-value and unadjusted p-value for each gene that passed the prescreening. These genes will be marked by different colors on the volcano plot, depending on the prescreening algorithm used. You can also hover over the markings on the volcano plot to see the specific gene name so that you can remove or add these genes from your BMD datasets on the next tab.

Note: If the unadjusted or adjusted p-value of a gene is too close to 0, the system may not include these on the plot because $\text{Log}(p\text{-value})$ will be undefined.

If you are not satisfied with the prescreening results, you can change the settings and re-execute the prescreening algorithms. Like before, an email will be sent to your inbox when the prescreening has completed.

Summary of Results:

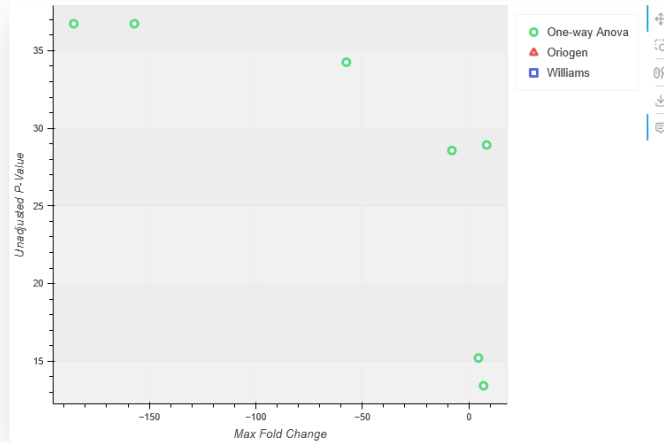
One-way Anova: 7 genes passed the screening

Williams Test: 7 genes passed the screening

Oriogen: 7 genes passed the screening

Volcano Plots

Max Fold vs Negative log of Unadjusted P-Value



Max Fold vs Negative log of Adjusted P-Value

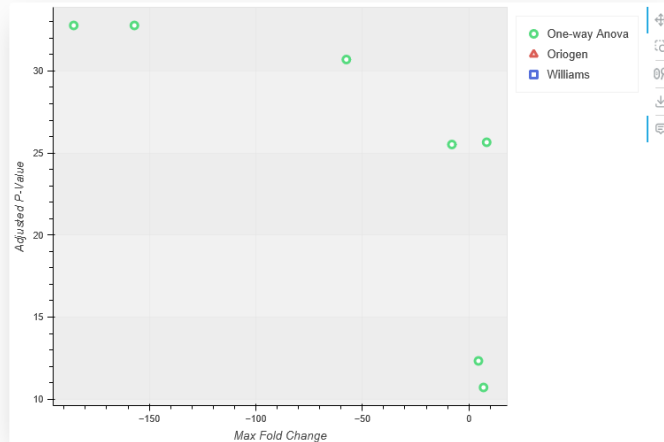


Figure 4.4. Preprocessing Summary and Volcano plots

iii. Explaining the algorithms used

Criteria including fold changes, P-values, adjusted P-values are applied to filter gene expression data.

a. Fold change

Fold change (f) is calculated using the following equation. The default value of f is 2, which means dose-response data with f smaller than 2 are filtered.

$$f = \begin{cases} \text{abs max} \left(-\frac{1}{\text{base}^{(x_i - x_0)}} \right); & \text{if } \text{base}^{(x_i - x_0)} < 1 \\ \text{abs max} \left(\text{base}^{(x_i - x_0)} \right); & \text{if } \text{base}^{(x_i - x_0)} \geq 1 \end{cases},$$

where x_i is the response at i -th dose level, x_0 is the response at control level, and base is the log base of data log transformation.

b. P-value and adjusted P-value

P-values are calculated by one way ANOVA and trend test (Williams test and Oriogen). With the P-values, adjusted P-values are calculated by Benjamini-Hochberg methods.

One-way Anova

One way ANOVA is a well-known test to determine whether there are any statistical differences between the means of the experiment group and the control group. The ANOVA produces an F-statistic, the ratio of the variance calculated among the means to the variance within the samples. A higher F-ratio implies that the samples were drawn from populations with different mean values. The F-ratio gives a P-value to determine whether significant differences exist in the experiment and control groups.

Williams Test and Oriogen

We adapt William's trend test (Williams 1971, 1972) and Oriogen (Peddada et al. 2005) to identify genes having a monotonical trend with respect to doses. That is, the maximum likelihood estimate (MLE) of mean response at i-th level is estimated by equation (1) and the test statistic (T) is calculated by equation (2). Permutation and bootstrap methods are applied to calculate the probability (P-value) that T_i (i-th permutation or bootstrap) is larger than T .

$$\hat{\mu}_i = \begin{cases} \max_{1 \leq u \leq i} \min_{1 \leq v \leq K} \frac{\sum_{j=u}^v n_j \bar{X}_j}{\sum_{j=u}^v n_j}; & \text{if increasing} \\ \min_{1 \leq u \leq i} \max_{1 \leq v \leq K} \frac{\sum_{j=u}^v n_j \bar{X}_j}{\sum_{j=u}^v n_j}; & \text{if decreasing} \end{cases}, \quad (1)$$

where $\hat{\mu}_i$ is the MLE of μ_i , K is dose level, i is the index of treatments group ($i = 1, \dots, K$), n_j is the number of samples at j-th level, and $\bar{X}_j$ is the mean response at j-th level.

$$T = \text{abs max} \left(\frac{\hat{\mu}_i - \bar{X}_0}{s \sqrt{\frac{1}{n_i} + \frac{1}{n_0}}} \right), \quad (2)$$

where $\bar{X}_0$ is the mean response at the control level, s is an unbiased estimate of within group standard deviation, n_i is the number of samples at i-th level, and n_0 is the number of samples at control level.

D. Data Preparation

After genes have been selected out with preprocessing, you have another opportunity to remove genes from your analysis in the data preparation step. This tab, shown in Figure 4.5, allows you to add and remove specific genes, and create multiple datasets for the BMD and Pathway analyses.

To create a dataset, follow the three steps below:

- (1) Give the dataset an identifiable name in the "Dataset Name" text box.
- (2) Choose the preprocessing algorithm which you will select genes from for the dataset.
- (3) Add the desired genes to the dataset using the checkbox in the gene's row. The sliders on the right-hand side of the page can be used to add all genes that fit those three parameters.

Even after adjusting the sliders, you can individually add or remove genes using the check box.

- (4) When you are satisfied with the dataset, click the green “Create” button on the bottom left side of the page.

You should see the dataset name appear on the left column of the page. If you want to edit this dataset, click on the name, and then change any of the settings previously chosen. After you are finished editing, save your changes by clicking the green “Update” button in the bottom left corner. To add an additional analysis, click the plus icon next to “Datasets” in the left column, then follow steps (1) – (4) again. You may include up to three datasets per genomic analysis. If you need to use more datasets, you can begin a new genomic analysis. If you want to delete a dataset, click the three dots icon next to the dataset name, then click “Delete”. This will remove that dataset from your list to be analyzed. If you had three datasets and then delete one, you can add another.

Bayesian BMD Dashboard Help About FAQ Contact

GENOMIC EXAMPLE

Progress: Data Input → Data Pre-processing → **Data Preparation** → Model Selection → BMD Settings → MCMC Settings → BMD Results → Pathway Preparation → Pathway Results → Share/Review

Datasets + Dataset Name: Dataset 1

Preprocessing Algorithms:
Origen, One-Way Anova, Williams

Probe ID	Gene Name	P-Value	Adjusted P-Value	Max Fold Change	Anova	Williams	Origen	Excluded
XPC_28078	XPC	2.5183e-7	4.4071e-6	4.53	☒	☒	☒	☐
EREG_21427	EREG	1.4890e-6	2.2336e-5	6.90	☒	☒	☒	☐
ACAA1_48	ACAA1	3.9124e-13	8.2161e-12	7.95	☒	☒	☒	☐
IL1B_3325	IL1B	2.7378e-13	7.1868e-12	8.31	☒	☒	☒	☐
ADH1B_20117	ADH1B	1.3323e-15	4.6629e-14	57.46	☒	☒	☒	☐
CYP2C8_15145	CYP2C8	1.1102e-16	5.8287e-15	156.98	☒	☒	☒	☐
FABP1_11582	FABP1	1.1102e-16	5.8287e-15	185.42	☒	☒	☒	☐

Create **Back** **Next**

Figure 4.5. “Dataset Preparation” tab for genomic BMD analysis

E. Model Selection

Before specifying BMD settings, you must select the models to be used in the model fitting part of the analysis.

The available dose-response models are:

- 1) Linear Model:

$$f(dose) = a + b \times dose, a > 0$$

- 2) Power Model:

$$f(dose) = a + b \times dose^g$$

$$a > 0, g \geq restriction$$

3) Hill Model:

$$f(dose) = a + \frac{b \times dose^g}{c^g + dose^g}$$

$$a > 0, c > 0, g \geq restriction$$

4) Exponential 2 Model:

$$f(dose) = a \times e^{b \times dose}, a > 0$$

5) Exponential 3 Model:

$$(dose) = a \times e^{b \times dose^g}$$

$$a > 0, g \geq restriction$$

6) Exponential 4 Model:

$$f(dose) = a \times (c - (c - 1) \times e^{-b \times dose})$$

$$a > 0, b > 0, c > 0$$

7) Exponential 5 Model:

$$f(dose) = a \times (c - (c - 1) \times e^{-(b \times dose)^g})$$

$$a > 0, b > 0, c > 0, g \geq restriction$$

To include a model in the analysis, click the checkbox on the right side of the model's row in the table displayed in the "Model Selection" tab (shown in Figure 4.6). When you check the box, the model is automatically added and saved to your analysis. Once you have added all the models you wish to include, press "Next" in the bottom right corner to begin specifying BMD settings.

[Dashboard](#)
[Help](#)
[About](#)
[FAQ](#)
[Contact](#)

GENOMIC EXAMPLE

Preview

DREAM Assist

Data Input
Data Pre-processing
Data Preparation
Model Selection
BMD Settings
MCMC Settings
BMD Results
Pathway Preparation
Pathway Results
Share/Review

Models

Select the models you want to include

Model Name	Formula	Include
Linear	$f(dose) = a + b \times dose$	☒
Power	$f(dose) = a + b \times dose^c$	☒
Hill	$f(dose) = a + \frac{b \times dose^c}{e^d + dose^c}$	☒
Exponential 2	$f(dose) = a \times e^{b \times dose}$	☒
Exponential 3	$f(dose) = a \times e^{b \times dose^c}$	☒
Exponential 4	$f(dose) = a \times [c - (c - 1) \times e^{-b \times dose}]$	☒
Exponential 5	$f(dose) = a \times [c - (c - 1) \times e^{-(b \times dose)^c}]$	☒

Back

Next

Figure 4.6. “Model Selection” tab for genomic BMD analysis

F. BMD Settings

i. Genomic BMD analysis steps

The next step in completing a genomic BMD analysis is to specify the BMD settings (the “BMD Settings” tab is shown in Figure 4.7). Two types of BMRs are available for genomic data: SD Change and Relative Change.

To create a BMD estimate to be analyzed, complete the following four steps:

- (1) Choose a name for the estimate. The default name for SD change is “BMD = *BMR Value* SD”. The default name for relative change is “BMD = *BMR Value* %”. If you would like to use a more identifiable name, enter this into the “Name” text box.
- (2) Select a BMR Type from the corresponding drop-down menu.
- (3) Specify the BMR value and enter this into the “BMR Value” text box.
- (4) Click “Save” at the bottom of the page to save these settings.

When the settings are successfully saved you should see the BMD settings name appear on the left column of the page. To edit these settings, click the name, then follow the previous steps to change any values. Once you have made the desired changes click the green “Update” button on the bottom of the page. If you want to instead cancel these changes, click the “Cancel” button on the bottom of the page. If you would like to add a new BMD setting for this analysis, click the plus icon next to “BMD” on the left side of the page. Three BMD settings are allowed for each genomic analysis. If you would like to delete this settings definition, click the setting’s name on the left side, and then click the red “Delete” button on the bottom of the page.

Figure 4.7. “BMD Settings” tab for genomic BMD analysis

ii. Explanation of the analysis calculation

Two options for defining the BMR values are provided based on the central tendency: a) relative change and b) standard deviation, which are described below. In BBMD, several BMRs can be defined.

$$f(BMD) \pm f(0) = \text{relative change} \times f(0),$$

$$f(BMD) \pm f(0) = k \times \text{standard deviation},$$

where $f(0)$ is the estimated response at zero dose, $f(BMD)$ is the response at BMD, relative change (e.g., 10%) and k (e.g., 1) are values defined by user, and standard deviation is estimated by models given observations. As a note, for every model, every posterior sample has an estimate for $f(0)$, standard deviation and therefore a BMD estimated value.

G. MCMC Settings

The final step before executing the BMD analysis is to specify the MCMC settings, the “MCMC Settings” tab is shown in Figure 4.8.

There are four different values that need to be specified in this tab. The first value is the number of Markov chain iterations, which can be 10,000 to 50,000 (inclusive) iterations per chain. Enter your value into the “Markov Chain Iterations” text box. Next, you need to specify the warmup percentage for each Markov Chain. This is the percentage of iterations discarded from the beginning of each chain; Therefore, those iterations will not be used for estimating model distributions. Put this percentage in the Warmup Percent (%) text box. The third value that needs to be specified is the number of Markov chains used in the analysis. Enter a number 1 to 3 (inclusive) into the “Number of Markov Chains” text box. Each chain will use the number of iterations previously specified. The final value is the random seed which is used for reproducing analysis results. The random seed can be 0 to 99,999 (inclusive). Enter this value in the “Random Seed text box”.

Once these values are specified, click “Save” to save the MCMC settings. If you are ready to execute the analysis, press “Execute” on the right side of the page. This will begin the BMD analysis execution. Once the execution has completed you will receive an email notification that your analysis is ready to be accessed. You will also receive a link which will send you back to the analysis to view the BMD results.

BBD Bayesian BMD

Dashboard Help About FAQ Contact

GENOMIC EXAMPLE

Preview DREAM Assist

Data Input Data Preprocessing Data Preparation Model Selection BMD Settings **MCMC Settings** BMD Results Pathway Preparation Pathway Results Share/Review

MCMC Settings

Markov Chain Monte Carlo (MCMC) settings using CmdStanPy

Markov chain iterations (per chain)

30000

Between 10,000–50,000, inclusive.

Warmup percent (%)

50

Between 10–90%, inclusive. This percent of iterations are discarded from the beginning of each chain, and not used for estimating the model distributions.

Number of Markov chains

1

Between 1–3, inclusive.

Random seed

91525

Between 0–99,999, inclusive.

Save

Back Execute

Figure 4.8. “MCMC Settings” tab for genomic BMD analysis

H. BMD Results

Once the BMD analysis has completed, the results will be displayed in the table on the “BMD Results” tab, shown in Figure 4.9. The names of the different datasets used will be shown on the left side of the page. Click the name of the dataset to view those specific results.

For each gene, the table displays the Model name, PPP value, model weight, and then the BMD, BMDU, and BMDL for each specified BMR. The BMR name will be displayed above the BMD value, and will be the name given when the BMR settings were created.

When you are finished reviewing the results, click “Next” in the bottom right corner to advance to the platform selection tab.

You can also skip the Pathway analysis and go to the “Share/Review” tab instead by clicking the name in the tab bar. This will allow you to export BMD results to an Excel spreadsheet immediately.

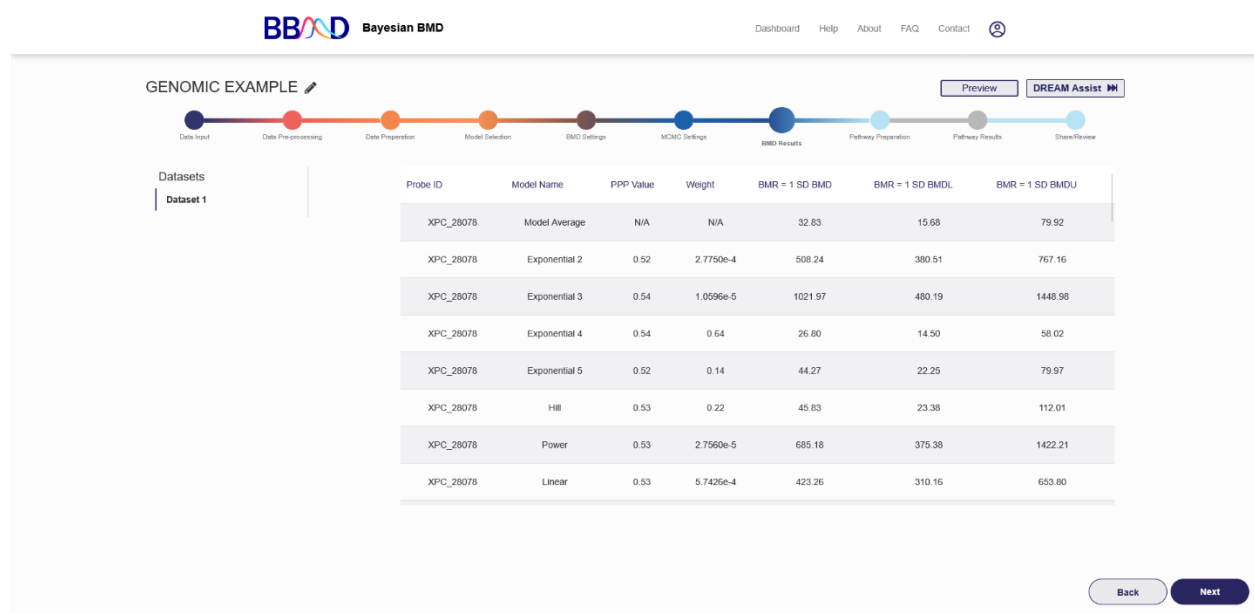


Figure 4.9. BMD Results after executing genomic analysis

I. Platform Selection

In this step, you must select the platform used for the pathway analysis from the drop-down menu. The pathway platforms that are available are shown in Figure 4.10. When you have chosen a platform, press “Execute” below the drop-down menu. Once the execution has completed, you will be able to advance to the next tab. Click “Next” in the bottom right corner to view the pathway results.

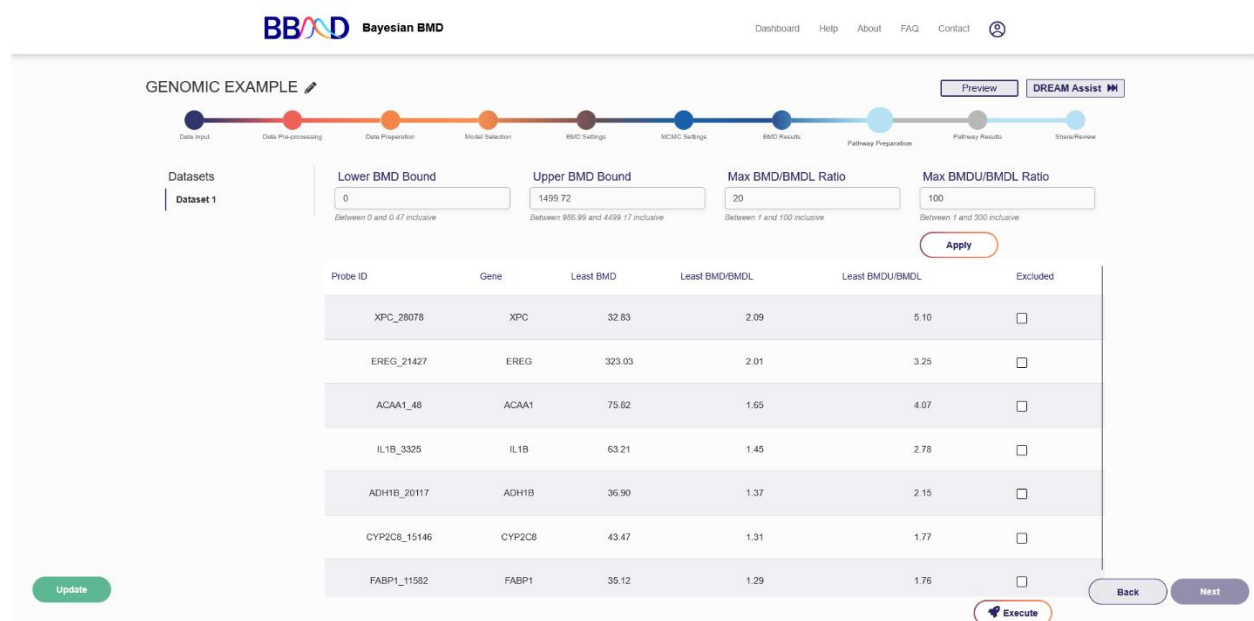


Figure 4.10. Pathway Preparation Screen

J. Pathway Analysis Results

i. How to view the pathway analysis results

The “Pathway Results” tab, shown in Figure 4.11, displays the results for the four different pathway types after completing a pathway analysis.

On the left side of the page, you can select the data set which you would like to view pathway results for. To switch datasets, click on the dataset name. The first section of this page is the pathway type selection. Click the pathway name to switch pathway types. The second section displays the platform summary. Summary information includes: the total number of pathways or total number of genes, number of unique genes, and then the Pathway BMD percentiles. The BMD definitions used for this analysis are the same BMDs that were used for the genomic BMD analysis.

Below the summary section the detailed Pathway BMD results are displayed in a large table. The columns for each table are slightly different, but each table always has a “Gene ID” column or “Pathway” column. If the text appears blue and underlined in the first column, there is a link to another resource to further analyze that pathway. Similarly, the last columns of the table are always the BMD results columns for each pathway or gene ID. If the text in the results cell is blue and underlined, clicking on it will bring up a plot showing this gene’s dose level with whiskers showing the BMDL and BMDU values. If any other text is blue and underlined in the table, hover over those cells because this row corresponds to multiple Probe IDs and Gene IDs.

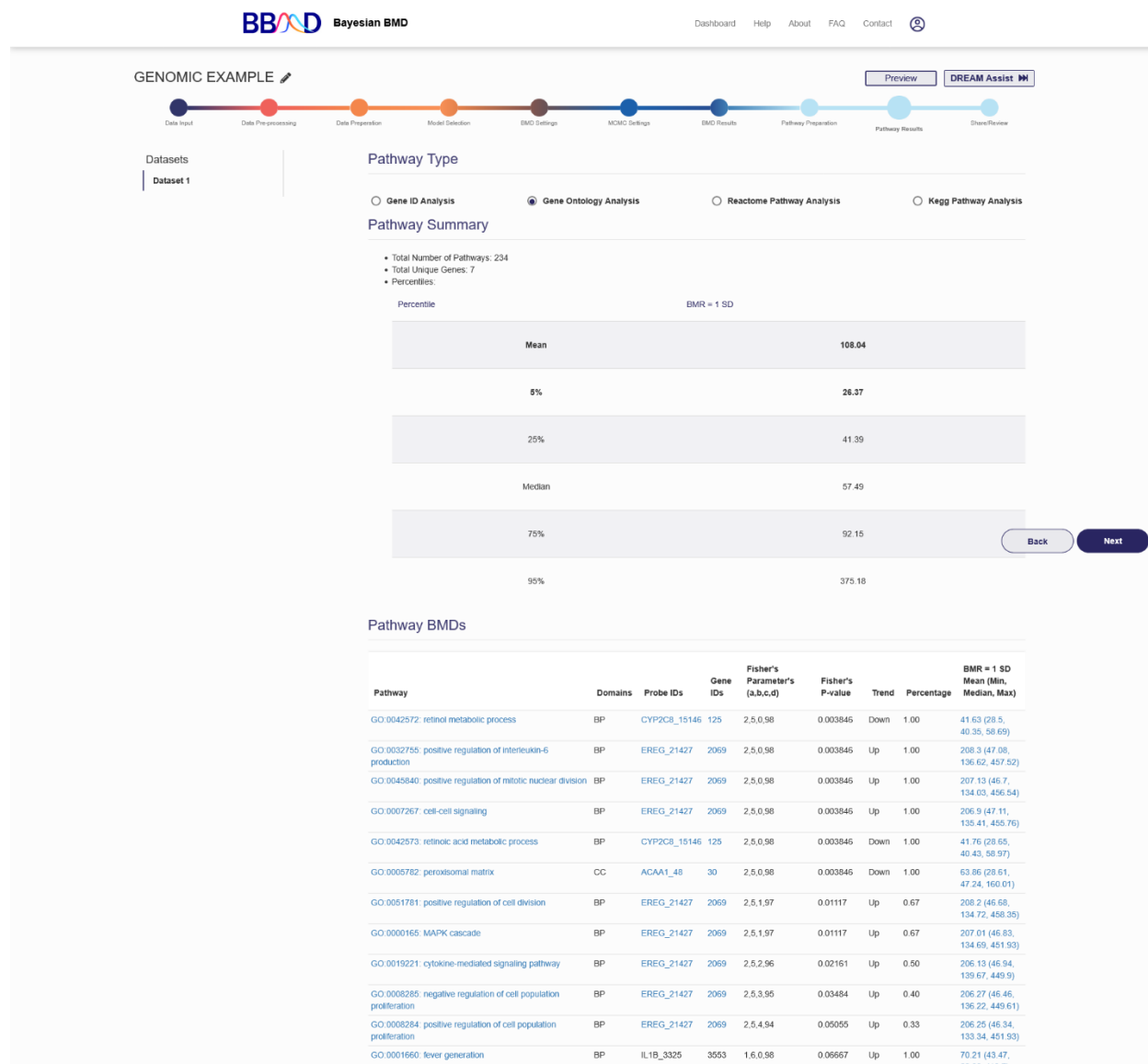


Figure 4.11. Pathway Analysis Results for "Gene Ontology Analysis"

ii. Algorithms used for pathway analysis

In the pathway analyses, platforms from Gene Expression Omnibus (GEO) are provided and users need to select one that associated with the uploaded genomic data. Four kinds of pathway analyses are provided to classify the BMA BMD analyses into significant pathways based on their NCBI Entrez Gene identifiers: a) Gene ID Analysis; b) GO Analysis (Mi et al. 2019); c) REACTOME (Fabregat et al. 2018) Pathway Analysis; and d) KEGG (Kanehisa and Goto 2000) pathway analysis. Each pathway analysis, the genes were matched to their associated categories, and the minimum, maximum, average, and median BMD were calculated for each category. For each analysis, a pathway summary table and a detailed pathway BMD table are displayed.

As for the four pathway methods, Gene ID Analysis simply translates the probe set identifiers to NCBI's Entrez Gene identifiers. GO Analysis utilizes 'go-basis.obo' and a python package

GOATOOLS (Klopfenstein et al. 2018) to group the Entrez Gene identifiers into three sub-ontologies: biological process, cellular component, and molecular function. The 'reactomepy' python module is used to access the REACTOME database, and API request is used to access the KEGG database.

For all the analyses, probe sets that measured more than one gene were removed from analyses. When different probe sets are associated with the same Entrez Gene identifiers, mean values of BMD are taken to represent the BMD value of the Entrez Gene identifiers. In order to determine whether the pathway is significant, P-values and percentages are calculated for each category. P-values are calculated based on Fisher's exact two-tailed test by comparing the numbers of genes with BMD estimates with the numbers of genes without BMD estimates. For each category, percentage is defined as the ratio of the number of genes with BMD estimates that are on this category to the total number of genes that are related to this category. The trend ('Up', 'Down' or 'Conflict') for each pathway is also provided. "Up" indicates > 60% genes in the category show up-regulation, "Down" indicates >60% show down-regulation and "Conflict" indicates neither "Up" or "Down" criteria were met.

K. Share/Review

The final tab for a genomic BMD analysis is the "Share/Review" page, shown in Figure 4.12. From this tab you can share this analysis with others and export the calculated results.

If you would like to share your analysis with others, you can change the settings from "Private" to "Send Back", "Share", or "Dream Assist". The "Share" setting allows you to send the created URL to others for them to access and review (but not edit) the analysis. The "Send Back" setting allows you to Return to the member that requested the DREAM Assist once reviewed. Lastly, the "DREAM Assist" setting will share this analysis with the DREAM Tech team, who can assist with the analysis. The "DREAM Assist" settings are coming soon.

You can also export the results of the analysis into Excel format. Before exporting results, you can customize the parts of the analysis included on the reports. By clicking "show more" on the right side of the page, the customization options will appear. From here you can select which datasets to include in the exported results. If you wish to change the report settings in the future, you can return to this analysis and export a new report. Exporting the results will send a link to your account's email where you can download the reports.

At any time during the updating or reviewing stage, if you want to change to another existing analysis, you can click the "Dashboard" button on the top right corner to switch to the summary page for the existing analyses and access another analysis.

BBMD

Bayesian BMD

Dashboard

Help

About

FAQ

Contact

GENOMIC EXAMPLE

Preview

DREAM Assist

Data Input

Data Pre-processing

Data Preparation

Model Selection

BMD Settings

MCMC Settings

BMD Results

Pathway Preparation

Pathway Results

Share/Review

Genomic Analysis Summary

Please review the Genomic Analysis Details and publication preferences. By selecting desired sections, you can either include or exclude tabs by selecting the check boxes below

Private

This analysis can only be viewed by you

Share

This analysis can be viewed by others

COMING SOON

DREAM Assist

Share your analysis with DREAM Tech and our team will assist with your analysis

Download

☐ Posterior Sample of BMD Estimates (Excel Format)

Show More

☐ Pathway BMD Estimates (Excel Format)

Show More

Back

Figure 4.12. “Share/Review” tab for a genomic BMD analysis

V. Probabilistic RfD Analysis

A. Introduction to the Probabilistic RfD Analysis

This analysis can be used to convert traditionally estimated BMD/BMDL, NOAEL/LOAEL to probabilistic reference dose given some additional input information.

B. Perform a Probabilistic RfD Analysis

An automatically generated name “New pRfD *Month Day Year, HH:MM AM/PM*”, is assigned to the newly started analysis. You can click the pencil button next to the analysis name, as seen in Figure 5.1, to make the name more identifiable.

This type of analysis has only two different tabs. The first tab is the data input tab where the analysis settings are specified, the other tab is the share tab.

To perform a Probabilistic RfD analysis, follow the 11 steps below on the “Data Input” tab:

- (1) Name the RfD analysis in the “RfD analysis name” box
- (2) Specify the dose units in the “Dose Units” box (This is not a required step)
- (3) (Optional) Choose a random seed for the “Random number seed” box
- (4) Input your prior BMD/BMDL or NOAEL/LOAEL data. If you choose to use a BMD/BMDL analysis as your point of departure selection, you must input the BMD value and BMDL value in the corresponding text boxes. If you choose to use a NOAEL analysis as the point of departure, along with the NOAEL, you must also choose an endpoint type from the drop-down list, and input NOAEL to BMD geometric mean and geometric standard deviation in the corresponding text boxes. Lastly, if you choose to use a LOAEL analysis, in addition to the same settings as a NOAEL analysis, you must also input the LOAEL to NOAEL geometric mean and geometric standard deviation in the corresponding text boxes.
- (5) Specify the Allometric Scaling settings. Select a test species from the “Test species” drop-down menu. Specify the test species body weight in the “Test species body weight” box. Choose a human body weight in the “Human body weight” box. Lastly, specify the Allometric scaling exponent mean and standard deviation in the two corresponding boxes.
- (6) Give an Animal to Human Uncertainty geometric mean and geometric standard deviation. These two values go in their respective boxes.
- (7) Specify Duration of Exposure details to extrapolate non-chronic exposures to chronic exposures. Select the duration of exposure from the drop-down menu. Then give the short-term exposure geometric mean and geometric standard deviation in the corresponding boxes.
- (8) (Optional) Add up to two additional uncertainties. Each additional uncertainty requires a name, geometric mean, and geometric standard deviation. These choices go in the respective boxes.
- (9) Specify the Human Variability geometric mean and geometric standard deviation in the corresponding boxes.
- (10) Give the target population-based incidence in the I* box.
- (11) Click the “Execute” button to execute the RfD analysis using the settings just specified.

When the analysis is completed, the name will be displayed on the left panel of the screen along with any other analyses. To edit the analysis settings, click the pencil icon on the right side of the page. You can also click the three dots next to the analysis name and choose “Edit RfD”. If you would like to delete this analysis, click the three dots and choose “Delete”. To add a new analysis, press the plus icon above the names of the existing analyses, and then follow steps (1) – (11) again.

BBMD Bayesian BMD Dashboard Help About FAQ Contact 23

Need More? With DREAM Premium, you'll get access to more features. Get the first 30 days free. [Get Started](#)

PRFD EXAMPLE

RfD **Create/update Probabilistic RfD**

All uncertainty factors are assumed lognormally distributed, so two parameters, (1) the geometric mean (gm) and (2) the geometric standard deviation (gpd) should be specified. In all cases, the gpd should be greater than or equal to 1. When an uncertainty factor is not applicable, a geometric mean and geometric standard deviation value of 1 and 1, respectively, are equivalent to an uncertainty factor that has no impact on the POD.

* indicates a required field.

RfD analysis name* **Dose units** **Random number seed* (for reproducibility)**

Point of departure selection

Please choose the Point of Departure (POD) type and value(s). Some POD types require additional uncertainty distributions to estimate a POD vector.

POD type* **BMDL*** **BMD***

Allometric Scaling

Test species* **Test species body weight (kg)*** **Human body weight (kg)***

Allometric scaling exponent (mean)* **Allometric scaling exponent (standard deviation)***

Animal to Human Uncertainty

Animal to Human Uncertainty (gm)* **Animal to Human Uncertainty (gpd)***

Duration of Exposure

Used to extrapolate non-chronic exposures to chronic exposures.

Duration of exposure* **Short term exposure uncertainty (gm)*** **Short term exposure uncertainty (gpd)***

Other uncertainties

Optional, added to allow a user to specify additional uncertainties.

Other uncertainty #1 (name) **Other uncertainty #1 (gm)** **Other uncertainty #1 (gpd)**

Other uncertainty #2 (name) **Other uncertainty #2 (gm)** **Other uncertainty #2 (gpd)**

Human variability

Human variability (gm)* **Human variability (gpd)***

Target population based incidence

* is target incidence level in human population, should be between 0.0001 and 0.5, where 0 is 0% of the population and 0.5 is 50% of the population.

I*

[Execute](#) [Cancel](#)

Figure 5.1. “Data Input” for a Probabilistic RfD analysis

The results, shown in Figure 5.2, will be displayed also be displayed on the “Data Input” tab. The first section in the results is a summary of the analysis settings. The next section is a table containing the input uncertainty/variability distributions. The next result displayed is the Point of departure (POD) distribution plot. The next plot is the HD₅₀, which is the estimated human

does at which 50% of the population has effects greater than or equal to the target magnitude of effect. The third plot is the HD_{M^1} , the estimated human dose where the population has 1% incidence of the target magnitude of effect, including inter-individual human variability. The next table contains the data in the plots for percentile and the corresponding POD, HD_{50} , and HD_{M^1} values. Below this table is the Probabilistic reference dose (RfD) value, defined as the 5th percentile HD_{M^1} , and the degree of uncertainty (90% CI). The next table is the relative contributions to HD_{M^1} uncertainty. The final plot shows the target incidence level I^* vs. Human dose.

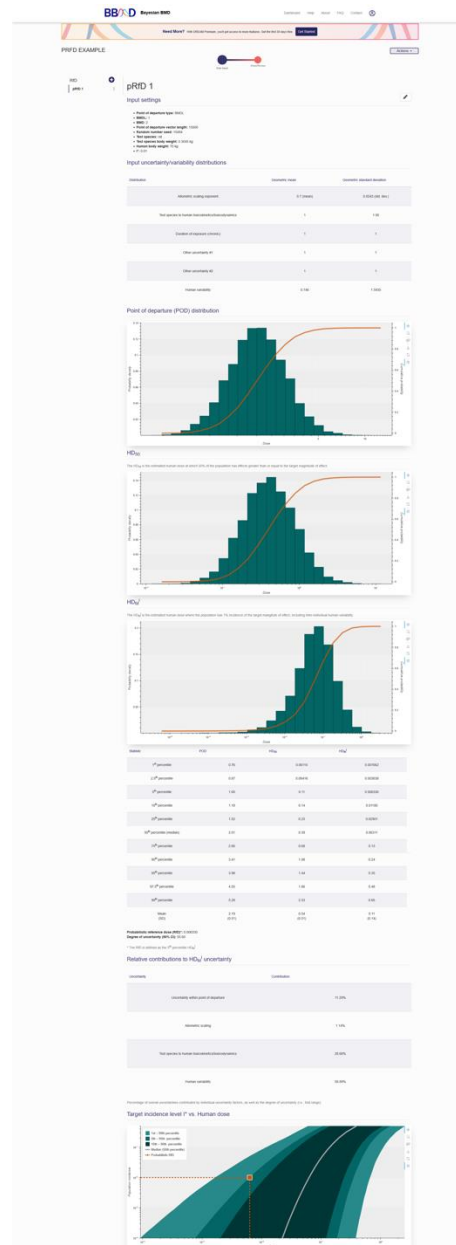


Figure 5.2. Results from a pRfD analysis

If you would like to share your analysis with others, you can change the settings from “Private” to “Send Back”, “Share”, or “Dream Assist”. The “Share” setting allows you to send the created

URL to others for them to access and review (but not edit) the analysis. The “Send Back” setting allows you to return to the member that requested the DREAM Assist once reviewed. Lastly, the “DREAM Assist” setting will share this analysis with the DREAM Tech team, who can assist with the analysis. The “DREAM Assist” settings are coming soon.

C. Interpret Probabilistic RfD Analysis Results

With the derived point of departure (POD) from an animal model and the user defined uncertainty factors, this POD is extrapolated to the equivalent human exposure (e.g., daily oral equivalent dose). Based on the dose response relationship of human extrapolation, the HD50 (50% of the population has effects greater than or equal to the target magnitude of effect) and HD_M¹ (estimated human dose where the population has 1% incidence of the target magnitude of effect, including inter-individual human variability) are derived on the RfD results page. The dose-population incidence curve below the summary table gives a visual expression of the statistic values in the HD_M¹ table. In addition, the contributions of uncertainty from each individual uncertainty factors are also summarized in the table on the RfD results page.

VI. BMD Analysis for Epidemiological Data

A. Introduction to BMD analysis for Epidemiological Data

This module is designed to perform BMD modeling and estimation for epidemiological data derived from case control studies.

When beginning a new analysis, an automatically generated name “New Epidemiological Analysis *Month Day Year, HH:MM PM*” is assigned to the analysis. You can click the pencil button next to the analysis name, as seen in Figure 6.1, to make the name more identifiable.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

EPIDEMIOLOGICAL EXAMPLE

Workflow: Data Input (selected) -> MCMC Settings -> Model Settings -> Execute Model Fit -> Model Fit -> BMD Estimates -> RD Estimates -> Risk at Dose -> Share/Review

Select Dataset Type

Select the type of data for your analysis in the dropdown below. Dataset type will determine the formatting of your data and cannot be changed once left this step.

Point Exposure

Odds Ratio or Relative Risk Percentile

95th Percentile

The percentile interval represented by the Odds Ratio Upper and Lower Bounds.

Dataset [Dose Case Non-Case OR/RR OR/RR_LOW OR/RR_HIGH]

	Dose	Case	Non-Case	OR/RR	OR/RR_LOW	OR/RR_HIGH
1						
2						
3						
4						
5						
6						
7						
8						

Load example

Input data with point exposure should have six columns, while input data with range exposure should have seven columns. For point exposure data, the first column should be the exposure level for each exposure group, and for range exposure data, the first and second columns should be the lower and upper bound of the exposure range of each exposure group. The remaining five columns for these two types of input data are the same, including (from left to right): number of cases, number of non-cases, median of OR/RR, lower limit of OR/RR, and upper limit of OR/RR.

Save Next

Figure 6.1. First page of a new epidemiological BMD analysis

B. Data input

i. Choosing an exposure type

The module for epidemiological data has two options for modeling exposure. You can either model your exposure as a point value or as a range. The range exposure function uses a bootstrap method to model the uncertainty in your exposures. This method requires more computational overhead and thus takes significantly longer to execute. You can choose which exposure modeling to use in the first drop down.

ii. How to input data into the system

To analyze an epidemiological dataset, you'll need to input it into the system. An epidemiological dataset has either 6 or 7 columns depending on which exposure type is selected. The first column is either the point exposure or dose or if range exposure is selected

there will be two columns one for the lower bound of the exposure and one for the upper bound. These columns are “Dose Low” and “Dose High”. If you’re using range exposure the Dose Low should be the lowest observed exposure in the group and Dose High should be the highest. Next, is the number of Cases then the number of Non-Cases. Next is the Odds Ratio followed by the lower and then upper bounds for the Odds Ratio. Whether you want to use the 95th or 90th percentile as the bounds for the odds ratio can be toggled using the drop-down menu. For an explanation on the Odds Ratio including how to calculate it see [here](#).

C. MCMC Settings

On this tab (shown in Figure 6.2), you can specify some settings for the MCMC algorithms.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

EPIDEMIOLOGICAL EXAMPLE Preview

Progress: Data Input → **MCMC Settings** → Model Settings → Execute Model Fit → Model Fit → BMD Estimates → RD Estimates → Risk at Dose → Share/Review

MCMC Settings

Markov Chain Iterations (per chain) 20000 Between 10,000–50,000, inclusive.	Warmup Percent (%) 50 Between 10–90%, inclusive. This percent of iterations are discarded from the beginning of each chain, and not used for estimating the model distributions.
Number of Markov Chains 1 Between 1–3, inclusive.	Random Seed 77584 Between 0–99,999, inclusive.

Back Next

Figure 6.2. MCMC settings

i. How to make and change settings

There are four different values that need to be specified in this tab. First, specify the number of Markov chain iterations, between 10,000 and 50,000 (inclusive) iterations per chain. Enter your value into the “Markov Chain Iterations” text box. Next, you need to specify the warmup percentage for each Markov Chain. This is the percentage of iterations discarded from the beginning of each chain; Therefore, those iterations will not be used for estimating model distributions. Put this percentage in the “Warmup Percent (%)” text box. Third, specify the number of Markov chains used in the analysis. Enter a number 1 to 3 (inclusive) into the “Number of Markov Chains” text box. Each chain will use the number of iterations previously specified. The final value is the random seed which is used for reproducing analysis results. The random seed can be 0 to 99,999 (inclusive). Enter this value in the “Random Seed text box”.

Once these values are specified, click “Next” to save the MCMC settings and move to ‘Model Settings’. Default settings are generally acceptable. However, results in the next step will provide important information that can help you judge if the MCMC settings are appropriate.

Based on our testing, **the default settings are adequate for most of the commonly seen dose-response shapes, so we suggest you use the default settings for your initial run.**

ii. How MCMC settings may impact the results

“Iterations” is the length of MCMC chain, i.e., the number of posterior samples in each MCMC chain. Default value is 30,000. The allowable range is any integer between 10,000 and 50,000.

“Number of chains” is the number of Markov Chains to be sampled. Default value is 1. Allowable range is 1 - 3.

“Warmup percent (%)”, the percent of sample in each Markov Chain will be discarded from the final posterior sample. Default value is 50% with an allowable range of 10% - 90%.

“Seed” is random seed number used in the MCMC algorithms. The number is randomly generated, but you can specify the number for the purpose of reproduction.

D. Model Settings

After the MCMC settings tab is the model settings tab. In this tab you choose which dose response models to fit to your dataset. The Epidemiological module has the same 8 models as the continuous module. Briefly these are:

Exponential 2
Exponential 3
Exponential 4
Exponential 5
Hill
Power
Michaelis Menten
Linear

You can choose individual models or choose “Standard Models” to add all 8 models. After choosing your models click execute to begin fitting the models.

E. Model Fit Results

On the “Model Fit Results” tab, the model fitting results obtained from the previous step are displayed. Click the name of one of the models on the left panel, then the results will be shown on the right (as shown in Figure 6.3) These results include the textual output of model parameter estimation, dynamic dose-response plot, posterior predictive p-value, model weight, correlation matrix, and graphical output of posterior sample of the model parameters (hidden by default). When click “Hide Parameters”, the parameter charts for each parameter in the model are displayed as shown in Figure 6.4.



Model Fit Summary

Seed used: 77584

Epidemiological Linear

Power

Epidemiological Linear fit summary

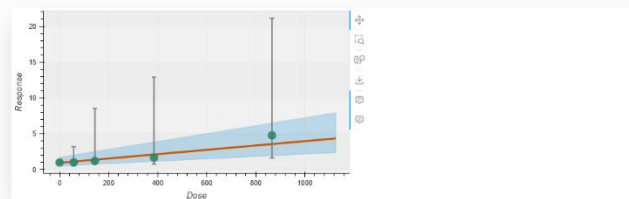
Model Formula

$$f(dose) = a + b \times dose$$

	Mean	HCSE	StdDev	2.5%	5%	50%	95%	97.5%	N_Eff	R_hat
a	0.954983	0.000112	0.000922	0.935528	0.938635	0.955001	0.971429	0.974492	7876	1.000006
b	2.631919	0.001126	0.180105	2.426730	2.454638	2.617250	2.786280	2.816478	8064	1.000258
tp	845.914400	0.017780	1.257648	842.580200	843.447000	846.255400	847.248900	847.318000	5003	1.000140
signa	0.954983	0.000067	0.000269	0.935320	0.935327	0.963448	0.973587	0.974810	8766	0.999954

Posterior predictive p -value for model fit: 0.477

Model weight: 8.2976e-31



[Back](#)

Next

Correlation Matrix

	a	b	sigma
a	1	-0.29	0.004285
b	-0.29	1	-0.006693
sigma	0.004285	-0.006693	1

Parameter Charts

Show Parameters

Figure 6.3. Results Shown on the “Model fit Results” Page

Parameter Charts

 Hide Parameters

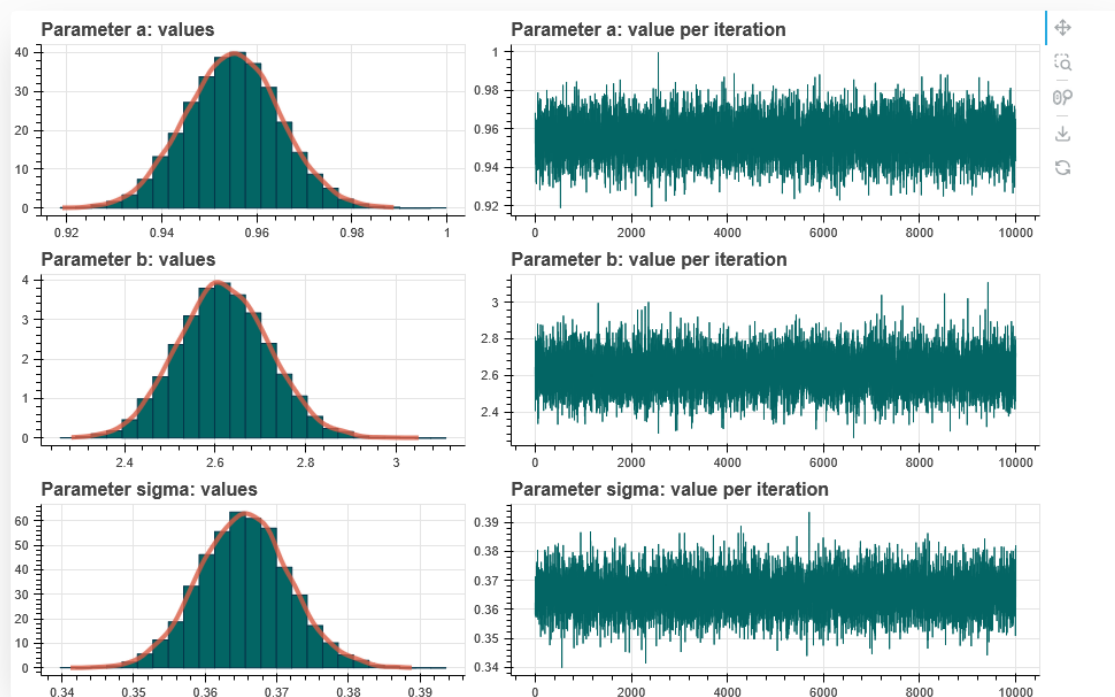


Figure 6.4. Parameter Charts

i. Parameter estimation results

The parameter estimation results, displayed in a table under the model formula, show the statistical summary for the estimated posterior distributions of parameters in the given dose-response model. These results are obtained directly from PyStan's fit output, including some important statistics for model parameters and diagnostic indicators for the MCMC algorithms. The mean, standard error of the mean (MCSE), standard deviation (StdDev), various quantiles (2.5%, 25%, 50%, 75%, and 97.5%), and quantities indicating effective sample size (N_Eff) and chain convergence (Rhat) for each model parameter derived from the posterior distribution of each parameter, as well as information regarding the MCMC execution are summarized in the table. As a note, the "Rhat" can be used to judge if the MCMC chains have converged properly. If the Rhat value is larger than 1.05, you may consider increasing the length of MCMC chains to get better convergence².

ii. Posterior Predictive P-Value

² Detailed explanation on the Stan outputs can be found at: <https://github.com/stan-dev/stan/releases/download/v2.9.0/stan-reference-2.9.0.pdf>

A posterior predictive p-value (PPP value) is reported below the dynamic dose-response plot. The PPP can be approximated by counting the predicted responses that satisfy the inequality out of the entire posterior sample space. This indicator can be used to judge if the fitting of this particular model is adequate. A large or small p-value means that a discrepancy in predicted data is very likely, further indicating a poor fit. Practically, if the PPP value is between 0.05 and 0.95, then the fitting is adequate. The calculation procedure of PPP value is briefly described below:

- (1) Use each bundle of parameters in the kept posterior sample to form a dose-response model and randomly generate case numbers, y^{rep} , at all dose levels in the original dataset
- (2) Use posterior sample of model parameters to calculate a test statistic for both the original data set (d, n, y) and the replicated data set (d, n, y^{rep}) . The test statistic used in this system is log-likelihood. For parameter values from l -th iteration, we have statistic $T(y, \theta^l)$ and $T(y^{rep}, \theta^l)$.
- (3) For $l = 1, \dots, L$ (the length of posterior sample), compare each pair of $T(y, \theta^l)$ and $T(y^{rep}, \theta^l)$, and count the number of $T(y, \theta^l) > T(y^{rep}, \theta^l)$, say M
- (4) The posterior predictive P-value is $\frac{M}{L}$

A detailed explanation on this procedure can be found in the Chapter of “Model checking and improvement” in *Bayesian Data Analysis* (Gelman et al).

iii. Posterior Model Weight

A model weight (w_j) for model j is calculated for each model included in the analysis as a statistic for cross-model comparison. The model weight was introduced by (Wasserman, 2000), using the following two equations. The value of each selected model j is calculated as follows:

where $\ln L_j$ is a loglikelihood value estimated using one set of posterior samples of model parameters of the j -th model, p_j is number of parameters in the j -th model, and n is the sample size in the data set.

When all models in the analysis have an equal prior weight, the posterior model weight of model j is calculated by m_j value estimated from model j divided by the sum of m values estimated from all models in the analysis as the following equation.

This function assumes equal model priors for all models selected, so the weight mainly indicates how well the model fits the data. To make the weight more reliable, we use 1000 sets of randomly selected posterior samples of model parameters to calculate the model weights. This model weights are further applied to the model averaged BMD calculation in the F. BMD Estimation section.

iv. Interactive Dose-Response Plot

A dynamic dose-response plot is shown below the text box. This plot includes original dose-response data and a fitted curve with its 90th percentile interval shaded in blue. When you

move your mouse over the dose-response curve, the estimated median and the 5th and 95th percentiles at a particular dose level will display. When you move your mouse over a data point from your inputted dataset, the dose, N, incidence, and the response percentile will also be displayed. Other information displayed in this figure includes the PyStan version, the lower bound placed on the power parameter (if applicable), the posterior predictive p-value (PPP value) for model fit and model weight for cross-model comparison.

v. Correlation Matrix

The fourth item displayed is the correlation matrix for the different model parameters. The correlation matrix is to show the correlation coefficients between different model parameters and is calculated using posterior samples.

vi. Plots for parameter posterior sample

If you click the “Show parameter charts”, two plots (posterior sample trace plot and estimated probability density plot) will be displayed for each of the parameters in this dose-response model.

Basically, this is the results display tab, meaning that you can only review the results, not give the system additional inputs to modify the results.

F. BMD Estimation

On this page, you can calculate the BMD estimates of your interest. The settings for epidemiological BMD estimates are similar to those for continuous data but differ in a few key ways. For Epidemiological BMDs each BMR is expressed in terms of a relative change from the response at some background exposure level. Figure 6.5 shows the screenshot for epidemiological BMD estimation

You can change the name of the BMD Settings using the first field. This has no effect on the actual BMD calculations, but does make it easier to navigate the BMD settings page when you have multiple BMDs. Next is the “Benchmark Response Value”. The BMR is expressed in terms of relative change from the background value. There are three options to specify the “Background Exposure Level”. Each option can be selected from the drop down menu. The first “Reference Group” sets the Background Exposure Level equal to the lowest exposure group. If you’re using an exposure range this value is the halfway between the lower and upper bounds of the range. The second background exposure option is to set the background exposure to zero. Finally, the “Custom” option allows the user to set the background exposure to any value. The system also allows you to specify the prior weight for each model. These prior weights should **not** consider the model fits on the previous tab as this information is already included in the algorithm to calculate the posterior weights. After you specify the settings click “Execute” to calculate the BMD.

BMD

+

Update BMD Settings

BMD setting name

Relative Change = 10%

Benchmark response value

0.1

Must be within the range of (0, 3.80).

Background Exposure Type

Reference Group

Background Exposure Level

0.12

Must be a non-negative value.

Figure 6.5. Epidemiological BMD settings

After executing the BMD will be displayed as a summary table and posterior plots of the BMD for the model average and for each model will be displayed. An example for epidemiological BMD estimation can be seen in figure 6.6 and Figure 6.7.

BMD

+

Relative Change = 10%

BMD summary table

Statistic	Model average	Exponential 2	Exponential 3	Exponential 4	Exponential 5	Hill	Power	Michaelis Menten	Linear
Prior model weight	N/A	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250	0.1250
Posterior model weight	N/A	0.0000	0.7709	0.0000	0.0249	0.0499	0.1543	0.0000	0.0000
Fraction with BMDs	100%	100%	100%	100%	100%	100%	100%	100%	100%
BMD (median)	110.87	54.24	105.75	30.49	153.07	142.41	139.83	30.42	31.75
BMDL (5 th percentile)	87.30	52.18	85.88	28.30	128.30	118.26	116.81	28.45	29.65
25 th percentile	100.08	53.38	97.38	29.54	142.23	132.22	129.77	29.63	30.87
Mean (SD)	114.84 (20.67)	54.27 (1.30)	106.55 (13.23)	30.52 (1.40)	154.42 (18.18)	142.78 (15.27)	140.45 (15.15)	30.48 (1.26)	31.79 (1.34)
75 th percentile	126.70	55.14	115.07	31.40	165.07	152.81	150.52	31.30	32.69
BMDU (95 th percentile)	154.99	56.44	129.41	32.92	183.57	168.85	166.03	32.62	34.04

Figure 6.6. Epidemiological BMD estimation summary table

BMD estimates

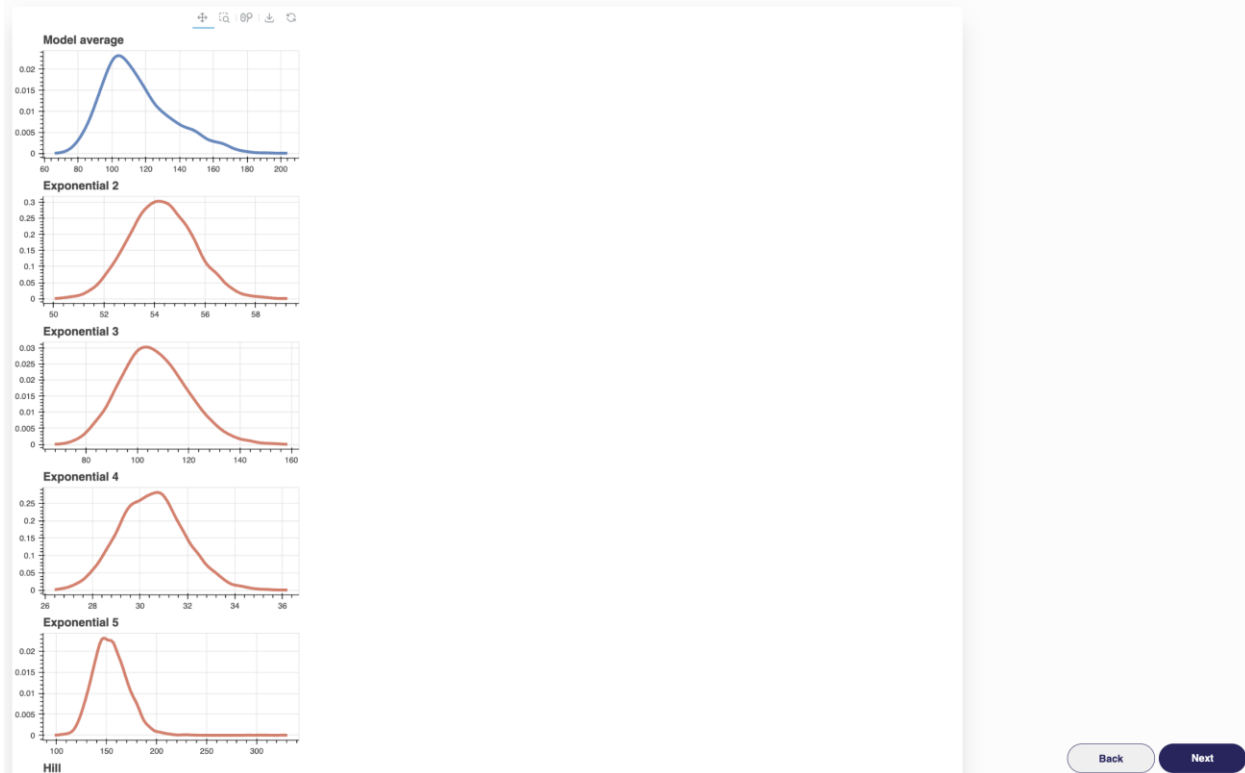


Figure 6.7. Posterior Plots of Epidemiological BMDs

VII. Multisite Tumor BMD Analysis

A. Introduction to Multisite Tumor BMD Analysis

This module is designed to perform BMD modeling and estimation for multiple tumor sites. Be aware this module makes the common assumption that each tumor site is conditionally independent (conditional on the dose level).

When beginning a new analysis, an automatically generated name “New MS Combo Run *Month Day Year, HH:MM AM/PM*” is assigned to the analysis. You can click the pencil button next to the analysis name, as seen in Figure 7.1, to make the name more identifiable.

BBMD Bayesian BMD Dashboard Help About FAQ Contact

MS COMBO EXAMPLE [Pencil Icon] [Preview]

Data Input MCMC Settings Model Settings Execute Model Fit Model Fit BMD Estimates Share

Tumor Sites

Add and name the sites of your tumors. Each tumor site should have a unique name. Please add your tumor sites before setting your data in the table below.

Tumor Site 1 [Delete]

Tumor Site 2 [Delete]

[Add Tumor Site]

Dataset [Dose N Tumor Site 1 Tumor Site 2]

	Dose	N	Tumor Site 1	Tumor Site 2
1				
2				
3				
4				
5				
6				
7				
8				

Input your data here. Each tumor site column should list the incidences for that tumor site. Each tumor site combined with the Dose and N columns should form a valid dichotomous dataset.

[Save] [Next]

Figure 7.1. First page of a new MS Combo BMD analysis

B. Data input

i. Setting your tumor sites

The Multisite Tumor module allows you to analyze up to five tumor sites. The default setup has two tumor sites but additional tumor sites can be added using the “Add Tumor Site” button. If you add too many sites you can delete some using the delete buttons to the right of the name fields. Tumor sites can be renamed by editing the text in the text box. This doesn’t affect the BMD calculations but does serve to help organize and track each tumor site. The names must be unique.

ii. How to input data into the system

To analyze a dataset, you'll need to input it into the system. A Multisite Tumor dataset is made up of several dichotomous datasets combined together. Like a dichotomous dataset the first two columns are the dose and the sample size (N). Then each Tumor site should have the incidence for that tumors at that particular site listed in its column. So if you have a control group with 50 animals and three with a tumor at Tumor Site 1 and four with a tumor at Tumor Site 2 the first row of your dataset would be 0 50 3 4. This is true regardless of whether an animal has tumors at sites 1 and 2 or not.

C. MCMC Settings

On this tab (shown in Figure 7.2), you can specify some settings for the MCMC algorithms.

BBMD Bayesian BMD

Dashboard Help About FAQ Contact

MS COMBO EXAMPLE

Preview

Data Input MCMC Settings Model Settings Execute Model Fit Model Fit BMD Estimates Share

MCMC Settings

Markov Chain Iterations (per chain)
30000
Between 10,000–50,000, inclusive.

Warmup Percent (%)
50
Between 10–90%, inclusive. This percent of iterations are discarded from the beginning of each chain, and not used for estimating the model distributions.

Number of Markov Chains
1
Between 1–3, inclusive.

Random Seed
36048
Between 0–99,999, inclusive.

Back Next

Figure 7.2. MCMC settings

i. How to make and change settings

There are four different values that need to be specified in this tab. First, specify the number of Markov chain iterations, between 10,000 and 50,000 (inclusive) iterations per chain. Enter your value into the “Markov Chain Iterations” text box. Next, you need to specify the warmup percentage for each Markov Chain. This is the percentage of iterations discarded from the beginning of each chain; Therefore, those iterations will not be used for estimating model distributions. Put this percentage in the “Warmup Percent (%)” text box. Third, specify the number of Markov chains used in the analysis. Enter a number 1 to 3 (inclusive) into the “Number of Markov Chains” text box. Each chain will use the number of iterations previously specified. The final value is the random seed which is used for reproducing analysis results. The random seed can be 0 to 99,999 (inclusive). Enter this value in the “Random Seed text box”.

Once these values are specified, click “Next” to save the MCMC settings and move to ‘Model Settings’. Default settings are generally acceptable. However, results in the next step will provide important information that can help you judge if the MCMC settings are appropriate.

Based on our testing, **the default settings are adequate for most of the commonly seen dose-response shapes, so we suggest you use the default settings for your initial run.**

ii. How MCMC settings may impact the results

“Iterations” is the length of MCMC chain, i.e., the number of posterior samples in each MCMC chain. Default value is 30,000. The allowable range is any integer between 10,000 and 50,000.

“Number of chains” is the number of Markov Chains to be sampled. Default value is 1. Allowable range is 1 - 3.

“Warmup percent (%)”, the percent of sample in each Markov Chain will be discarded from the final posterior sample. Default value is 50% with an allowable range of 10% - 90%.

“Seed” is random seed number used in the MCMC algorithms. The number is randomly generated, but you can specify the number for the purpose of reproduction.

D. Model Settings

After the MCMC settings tab is the model settings tab. In this tab (Figure 7.3) you choose which dose response models to fit to your dataset. Due to the assumptions made about conditional independence only two of our 8 dichotomous models are suitable for the Multisite Tumor module. They are the Quantal Linear and Multistage models. Each model can be included by checking the associated box or excluded by unchecking the box. You can choose whether to use our objective priors or our empirical informative priors from the prior settings menu. At least one model must be selected before clicking “Execute”.

Bayesian BMD

Dashboard Help About FAQ Contact

MS COMBO EXAMPLE

Data Input MCMC Settings **Model Settings** Execute Model Fit Model Fit BMD Estimates Share

Model Settings

Models

☒ Quantal Linear Model

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose})$$

☒ Multistage (2nd order) Model

$$f(dose) = a + (1 - a) \times (1 - e^{-b \times dose - c \times dose^2})$$

Prior Settings

☒ Non-informative Prior (Objective Prior)

Use a fixed/unchangeable non-informative prior to model the dose-response data objectively.

☐ Empirical Informative Prior

Use predetermined distributions derived from the general database as an informative prior for dose-response modeling.

Back Execute

Figure 7.3. MS Tumor Model Settings

E. Model Fit Results

On the “Model Fit Results” tab, the model fitting results obtained from the previous step are displayed. Click the name of one of the models on the left panel, then the results will be shown on the right (as shown in Figure 7.4) These results include the textual output of model parameter estimation, the combined model formula, model weight, and dynamic dose response plot for each tumor site model (Figure 7.5).

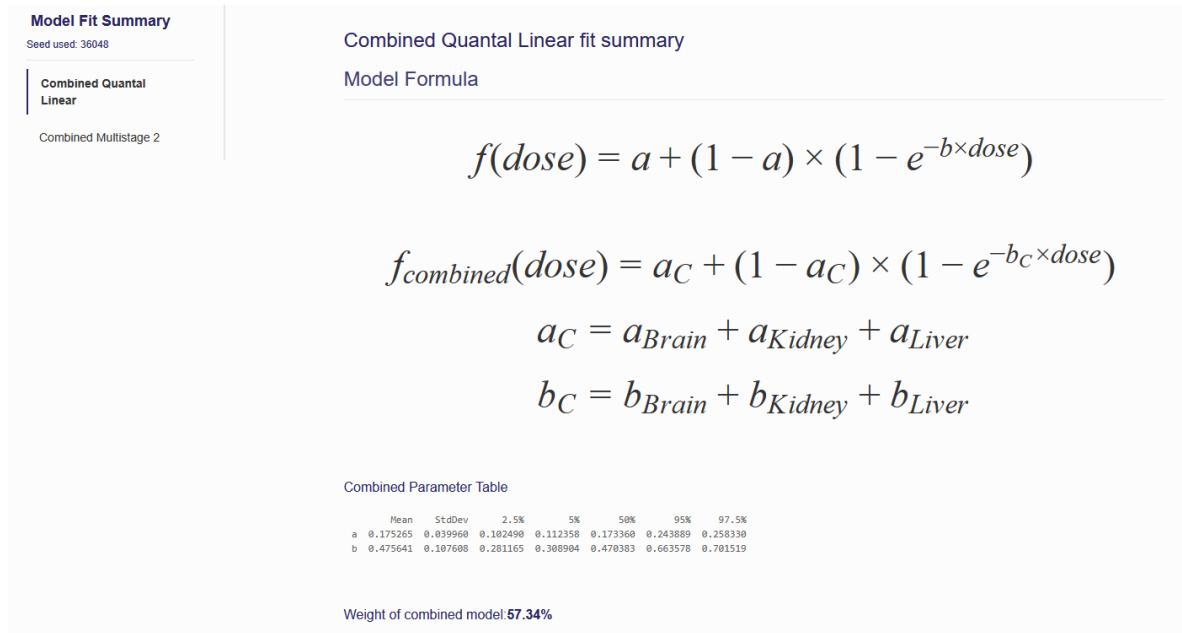


Figure 7.4. Results Shown on the “Model fit Results” Page

Individual Site Models

Brain

	Mean	MCSE	StdDev	2.5%	5%	50%	95%	97.5%	N_Eff	R_hat
a	0.025706	0.000203	0.017415	0.003418	0.004915	0.022001	0.059738	0.068816	7362	0.999999
b	0.299315	0.000836	0.070907	0.169959	0.188276	0.296650	0.422971	0.447088	7193	1.000000
1p__	-13.645100	0.014597	1.040810	-16.509400	-15.763700	-13.339200	-12.637200	-12.608800	5084	0.999941

Posterior predictive p -value for model fit: **0.647**

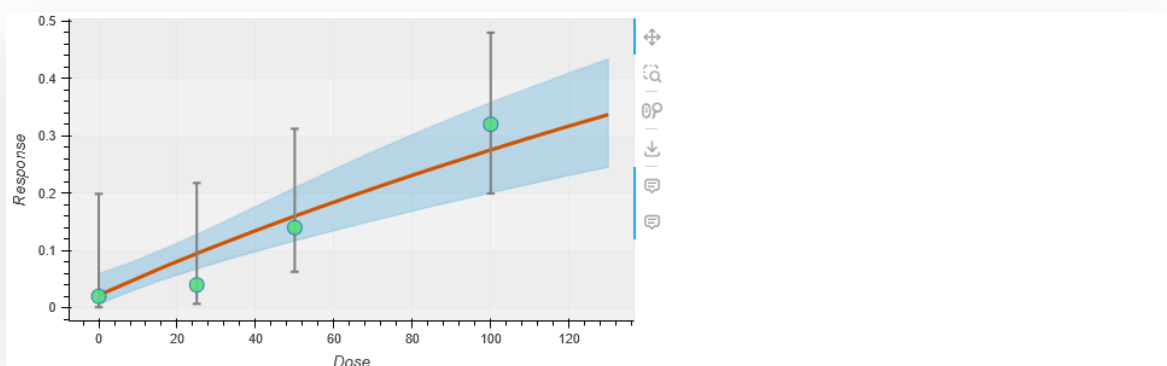


Figure 7.5. Example of an individual tumor site model

i. Parameter estimation results

The parameter estimation results, displayed in a table under the model formula, show the statistical summary for the estimated posterior distributions of parameters in the given dose-response model. The mean, standard error of the mean (MCSE), standard deviation (StdDev), and various quantiles (2.5%, 25%, 50%, 75%, and 97.5%) for each model parameter derived from the posterior distribution of each parameter are summarized in the table.

ii. Posterior Model Weight

A model weight (w_j) for model j is calculated for each model included in the analysis as a statistic for cross-model comparison. The model weight was introduced by (Wasserman, 2000), using the following two equations. The value of each selected model j is calculated as follows:

where ℓ_j is a loglikelihood value estimated using one set of posterior samples of model parameters of the j -th model, p_j is number of parameters in the j -th model, and n is the sample size in the data set.

When all models in the analysis have an equal prior weight, the posterior model weight of model j is calculated by m_j value estimated from model j divided by the sum of m values estimated from all models in the analysis as the following equation.

This function assumes equal model priors for all models selected, so the weight mainly indicates how well the model fits the data. To make the weight more reliable, we use 1000 sets of

randomly selected posterior samples of model parameters to calculate the model weights. This model weights are further applied to the model averaged BMD calculation in the F. BMD Estimation section.

iii. Individual Tumor Site Models

The models for each individual tumor site are shown under the combined model. Each individual model has its own textual output table, interactive dose response plot, and posterior predictive p value displayed.

F. BMD Estimation

On this page, you can calculate the BMD estimates of your interest. The Multisite Tumor BMD calculations are similar to those for dichotomous data.

You can change the name of the BMD Settings using the first field. This has no effect on the actual BMD calculations but does make it easier to navigate the BMD settings page when you have multiple BMDs. Next is the Benchmark Response Value. The BMR is calculated using both the added and extra risk definitions.

BMD

BMR = 10%

Update BMD Settings

BMD setting name

BMR = 10%

Benchmark response value

0.1

Must be within the range of (0, 1.00).

Execute Cancel Delete

Figure 7.6. MS Tumor BMD settings

After executing the BMD will be displayed as a summary table and posterior plots of the BMD for the model average and for each model will be displayed.

VIII. References

- Crump KS. 1995. Calculation of benchmark doses from continuous data. *Risk Analysis*, 15(1): 79-89.
- Gelman, A., Carlin, J.B., Stern, H.S., Rubin, D.B. (2003). *Bayesian Data Analysis*. Second Edition. Boca Raton, FL: Chapman and Hall/CRC Press.
- Slob W. 2002. Dose-response modeling of continuous endpoints. *Toxicological Science*, 66(2): 298-312.
- Wasserman, L. (2000). Bayesian model selection and model averaging. *Journal of Mathematical Psychology*, 44(1), 92-107.
- Williams D. 1971. A test for differences between treatment means when several dose levels are compared with a zero dose control. *Biometrics*:103-117.
- Williams D. 1972. The comparison of several dose levels with a zero dose control. *Biometrics*:519-531.
- Peddada S, Harris S, Zajd J, Harvey E. 2005. Oriogen: Order restricted inference for ordered gene expression data. *Bioinformatics* 21:3933-3934.
- Mi H, Muruganujan A, Ebert D, Huang X, Thomas PD. 2019. Panther version 14: More genomes, a new panther go-slim and improvements in enrichment analysis tools. *Nucleic acids research* 47:D419-D426.
- Fabregat A, Jupe S, Matthews L, Sidiropoulos K, Gillespie M, Garapati P, et al. 2018. The reactome pathway knowledgebase. *Nucleic acids research* 46:D649-D655.
- Kanehisa M, Goto S. 2000. Kegg: Kyoto encyclopedia of genes and genomes. *Nucleic acids research* 28:27-30.