



Data Mining & Exploration Program



Dipartimento di Scienze Fisiche
Università di Napoli "Federico II"



ISTITUTO NAZIONALE di ASTROFISICA
OSSERVATORIO ASTRONOMICO di CAPODIMONTE



CALTECH



Graphical User Interface User Manual

DAME-MAN-NA-0010

Issue: 1.4
Date: March 16, 2015
Author: M. Brescia, S. Cavuoti

Doc. : GUI_UserManual_DAME-MAN-NA-0010-Rel1.4





Data Mining & Exploration Program

DAME Program
“we make science discovery happen”

INDEX

1	Introduction	5
2	Purpose	5
3	GUI Overview	6
3.1	User Registration and Access	7
3.2	The command icons	10
3.3	Workspace Management	12
3.4	Header Area	13
3.5	Data Management	15
3.5.1	Upload user data	16
3.5.2	How to Create dataset files	17
3.5.2.1	Feature Selection	18
3.5.2.2	Column Ordering	19
3.5.2.3	Sort Rows by Column	20
3.5.2.4	Column Shuffle	21
3.5.2.5	Row Shuffle	22
3.5.2.6	Split by Rows	23
3.5.2.7	Dataset Scale	23
3.5.2.8	Single Column Scale	24
3.5.3	Download data	24
3.5.4	Moving data files	24
3.6	Plotting and Visualization	25
3.6.1	Plotting	25
3.6.2	Visualization	31
3.7	Experiment Management	31
3.7.1	Re-use of already trained networks	36

TABLE INDEX

Tab. 1 – Header Area Menu Options	14
Tab. 2 – Abbreviations and acronyms	40
Tab. 3 – Reference Documents	41
Tab. 4 – Applicable Documents	42



Data Mining & Exploration Program

FIGURE INDEX

Fig. 1 – Suite functional hierarchy.....	6
Fig. 2 –The user registration/login form to access at the web application.....	8
Fig. 3 – The user registration form	9
Fig. 4 – An example of e-mail received by the user after submission of registration info	9
Fig. 5 – The Web Application starting main page (Resource Manager).....	10
Fig. 6 – The Web Application main areas and commands.....	11
Fig. 7 – The right sequence to configure and execute an experiment workflow	12
Fig. 8 – the button “New Workspace” at left corner of workspace manager window.....	12
Fig. 9 – the form field that appears after pressing the “New Workspace” button.....	13
Fig. 10 –the active workspace created in the Workspace List Area.....	13
Fig. 11 –The GUI Header Area with all submenus open	14
Fig. 12 – The Upload data feature open in a new tab.....	16
Fig. 13 – The Upload data from external URI feature.....	16
Fig. 14 – The Upload data from Hard Disk feature	17
Fig. 15 – The Uploaded data file in the File Manager sub window.....	17
Fig. 16 – The dataset editor tab with the list of available operations.....	18
Fig. 17 – The Feature Selection operation – select columns and put saving name	19
Fig. 18 –The Feature Selection operation – the new file created.....	19
Fig. 19 –The Column Ordering operation – the starting view.....	19
Fig. 20 –The Column Ordering operation – new order to columns.....	20
Fig. 21 – The Column Ordering operation – new file created.....	20
Fig. 22 –The Sort Rows by Column operation – step 1.....	20
Fig. 23 –The Sort Rows by Column operation – step 2.....	21
Fig. 24 –The Sort Rows by Column operation – the new file created.....	21
Fig. 25 –The Column Shuffle operation – step 1	21
Fig. 26 –The Column Shuffle operation – the new file created.....	22
Fig. 27 –The Row Shuffle operation – step 1	22
Fig. 28 –The Row Shuffle operation – the new file created.....	22
Fig. 29 –The Split by Rows operation – step 1.....	23
Fig. 30 –The Split by Rows operation – the new files created	23
Fig. 31 –The Dataset Scale operation – step 1.....	24
Fig. 32 –The Single Column Scale operation – step 1	24
Fig. 33 – The Histogram tab	25
Fig. 34 – A multi layer histogram plot	26
Fig. 35 – The Scatter 2D tab	27
Fig. 36 – A multi layer scatter 2D plot.....	28
Fig. 37 – The Scatter Plot 3D tab.....	28
Fig. 38 – The Line Plot tab.....	29
Fig. 39 – A multi layer line plot.....	30
Fig. 40 – The visualization tab	31
Fig. 41 – Creating a new experiment (by selecting icon “Experiment” in the workspace)	32
Fig. 42 – The new tab open after creation of a new experiment with the list of available options	32
Fig. 43 – The new state of the experiment configuration tab after the selection of the model	33
Fig. 44 – The configuration options in the Train use case.....	34
Fig. 45 – The configuration options in the Test use case	34
Fig. 46 – The configuration options in the Run use case	34
Fig. 47 – The configuration options in the Full use case	35
Fig. 48 – Example of a web page automatically open after the click on the help button.....	35
Fig. 49 – Some different state of two concurrent experiments	35
Fig. 50 – An example of Classification_MLP training case for the XOR problem	36



DAta Mining & Exploration Program

<i>Fig. 51 – The popup status at the end of the XOR problem experiment.....</i>	<i>36</i>
<i>Fig. 52 – The list of output files after the XOR problem training experiment.....</i>	<i>37</i>
<i>Fig. 53 – The training error scatter plot mlp_TRAIN_errorPlot.jpg downloaded from the experiment output list (x-axis is the training cycle, y-axis is the training mean square error).....</i>	<i>37</i>
<i>Fig. 54 – The operation to “move” the trained network file in the Workspace input file list.....</i>	<i>38</i>
<i>Fig. 55 –the configuration for the Run use case in the XOR problem.....</i>	<i>38</i>
<i>Fig. 56 – the output of the TEST use case experiment in the XOR problem.....</i>	<i>39</i>



DAta Mining & Exploration Program

1 Introduction

The present document is part of the DAMEWARE Web Application Suite user-side documentation package. The final release arises from the very primer version of the web application (α release) which has been made available to public domain since July 2010. Currently this release has been more updated, by fixing residual bugs and by adding more functionalities and models.

The final developing team has spent much efforts to fix bugs, satisfy testing user requirements, suggestions and to improve the application features, by integrating several other data mining models, always coming from machine learning theory, which have been scientifically validated by applying them offline in several practical astrophysical cases (photometric redshifts, quasar candidate selection, globular cluster search, transient discovery etc). All cases dealing with time domain data rich astronomy. In this scenario, the α release has covered the role of an advanced prototype, useful to evaluate, tune and improve main features of the web application, basically in terms of:

- User friendliness: by taking care of the impact on new users, not necessarily expert in data mining or skilled in machine learning methodologies, by paying particular attention to the easiness of navigation through GUI options and to the learning speed in terms of experiment selection, preprocessing, setup and execution;
- Data I/O handling: easiness to upload/download data files, to edit and configure datasets from original data files and/or archives;
- Workspace handling: the capacity to create different work spaces, depending on the experiment type and data mining model choice;

Of course in the new release it was impossible to match all important and valid suggestions came from the α release testers. In principle not for bad will of developers, but mostly because in some cases, the requests would needed drastic re-engineering of some infrastructure components or simply because they went against our design requirements, issued at the very beginning of the project. Of course, this not implies necessarily that in next releases of the application these requests will not be taken into account.

Anyway, we tried to satisfy as much as possible main requests concerning the improvement of ease to use. Also in terms of examples and guided tours in using the available models. Don't forget that neophyte users should spend a certain amount of time to read this and other manuals to learn their capabilities and usability topics before to move inside the application. This is particularly true in order to understand how to identify the right association of functionality domain and the data mining model to be applied to your own science case. But we recall that this is fully reachable by gaining experience with time and through several trial-and-error sessions.

2 Purpose

This manual is mainly dedicated to drive users through the GUI options and features. In other words to show how to navigate and to interact with the application interface in order to create working spaces, experiments, to upload/download and edit data files. We will stop our discussion here at level of configuration of the models, for which specific manuals are available. This in order to separate the use of the GUI from the theoretical implications related to the setup and use of available data mining models.

The access gateway, its complete documentation package and other resources is at the following address:

<http://dame.dsf.unina.it/dameware.html>



Data Mining & Exploration Program

Last pages of this document host tables with “Abbreviations & Acronyms”, “Reference” and “Applicable” document lists and the acknowledgments. All over the document the references are labeled as [Rxx] for “Reference” documents and [Axx] for “Applicable” documents (xx is the incremental index as reported in the list tables). “Applicable” documents are not public references (technical documents internal to the DAME working group) included for quick technical references. Users external to the working group may ask to consult (privately) these documents by e-mail, motivating the reasons. The complete list of the internal documentation is available at the following address of the program official website: http://dame.dsf.unina.it/DAME_DOCUMENTATION_LIST.html.

3 GUI Overview

Main philosophy behind the interaction between user and the DMS (Data Mining Suite) is the following.

The DMS is a web application, accessible through a simple web browser. It is structurally organized under the form of working sessions (hereinafter named workspaces) that the user can create, modify and erase. You can imagine the entire DMS as a container of services, hierarchically structured as in Fig. 1. The user can create as many workspaces as desired and populate them with uploaded data files and with experiments (created and configured by using the Suite). Each workspace is enveloping a list of data files and experiments, the latter defined by the combination between a functionality domain and a series (one at least) of data mining models. From these considerations, it is obvious that a workspace makes sense if at least one data file is uploaded into. **So far, the first two actions, after logged in, are, respectively, to create a new workspace (by assigning it a name) and to populate it by uploading at least one data file, to be used as input for future experiments. The data file types allowed by the DMS are reported in the next sections.**

In principle there should be many experiments belonging to a single workspace, made by fixing the functional domain and by slightly different variants of a model setup and configuration or by varying the associated models.

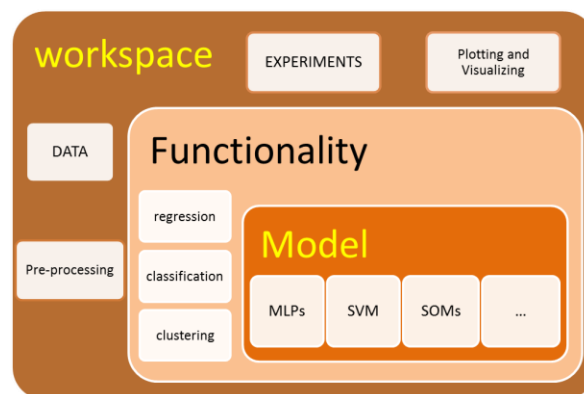


Fig. 1 – Suite functional hierarchy

By this way, as usual in data mining, the knowledge discovery process should basically consist of several experiments belonging to a specified functionality domain, in order to find the model, parameter configuration and dataset (parameter space) choices that give the best results (in terms of performance and reliability). The following sections describes in detail the practical use of the DMS from the end user point of view. Moreover, the DMS has been designed to build and execute a typical complete scientific pipeline



DAta Mining & Exploration Program

(hereinafter named workflow) making use of machine learning models. This specification is crucial to understand the right way to build and configure data mining experiment with DMS.

In fact, machine learning algorithms (hereinafter named models) need always a pre-run stage, usually defined as training (or learning phase) and are basically divided into two categories: supervised and unsupervised models, depending, respectively, if they make use of a BoK (Base of Knowledge), i.e. couples input-target for each datum, to perform training or not (for more details about the concept of training data, see section 3.5 below).

So far, any scientific workflow must take into account the training phase inside its operation sequence.

Apart from the training step, a complete scientific workflow always includes a well-defined sequence of steps, including pre-processing (or equivalently preparation of data), training, validation, run, and in some cases post-processing.

The DMS permits to perform a complete workflow, having the following features:

- A workspace to envelope all input/output resources of the workflow;
- A dataset editor, provided with a series of pre-processing functionalities to edit and manipulate the raw data uploaded by the user in the active workspace (see section 3.5 for details);
- The possibility to copy output files of an experiment in the workspace to be arranged as input dataset for subsequent execution (the output of training phase should become the input for the validate/run phase of the same experiment);
- An experiment setup toolset, to select functionality domain and machine learning models to be configured and executed;
- Functions to visualize graphics and text results from experiment output;
- A plugin-based toolkit to extend DMS functionalities and models with user own applications;

3.1 User Registration and Access

The DMS makes use (embedded to the end user) of the Cloud computing infrastructure, made by single PCs in combination with GRID resources. This requires a reliable level of security in order to launch jobs (experiments) in a safe and coordinated way. This level of security is obtained by an accounting procedure that foresees an initial registration for new users, in order to activate their account on the DAME Suite. After activation, all subsequent accesses will require login and password, as defined by the user at the registration stage. The user registration/login entry page is shown in Fig. 2.



Data Mining & Exploration Program

The screenshot shows the DAME web application interface. At the top is a blue banner with the text "DAta Mining & Exploration" and a "DAME" logo. Below the banner, on the left, is a description of DAME as an innovative, general purpose, Web-based, distributed data mining infrastructure. It also provides a link for news, documentation, and FAQ information. In the center, there is a "Technical Support" section with an email address and a "Skype helpdesk" button. On the right, there is a "Sign in" section with fields for "User Mail" and "Password", a "Login" button, and a "New on DAME WebApp?" section with a "Register Now" button. Below the "Register Now" button, there is a list of four steps for the registration procedure. At the bottom, there is a row of logos for various organizations, including INAF-OACN, INDA, and EUROVA.

DAta Mining & Exploration

DAME

DAME (Data Mining & Exploration) is an innovative, general purpose, Web-based, distributed data mining infrastructure specialized in Massive Data Sets exploration with machine learning methods.

For news, documentation and FAQ information please click [HERE](#)

Technical Support
helpdame AT gmail.com

Skype helpdesk
[Call me!](#)

Sign in

User Mail:

Password:

Login

New on DAME WebApp?

Register Now

You can obtain the access by following a simple registration procedure:

1. Compile the registration form (click *Register Now* button);
2. Immediately after you will receive by e-mail a welcome message;
3. Check for an e-mail message with your account confirmation;
4. Go back at this page and sign in;

Logos: INAF-OACN, INDA, EUROVA, etc.

Fig. 2 –The user registration/login form to access at the web application

New users must be registered by following a very simple procedure requiring to select “Register Now” button on that page.

The registration form requires the following information to be filled in by the user (all fields are required):

- Name of the user;
- Family name of the user;
- User e-mail: the user e-mail (it will become his access login). It is important to define a real address, because it will be also used by the DMS for communications, feedbacks and activation instructions;
- Country: country of the user;
- Affiliation: the institute/academy/society of the user;
- Password: a safe password (at least 6 chars), without spaces and special chars;



Data Mining & Exploration Program

Fig. 3 – The user registration form

After submission, an e-mail will be immediately sent at the defined address (Fig. 4), confirming the correct coming up of the activation procedure.

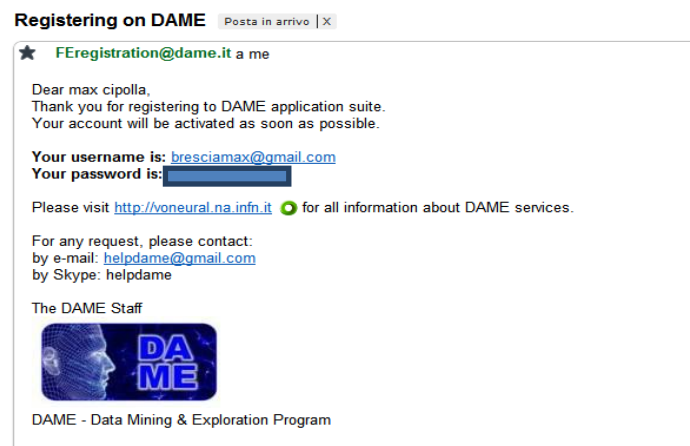


Fig. 4 – An example of e-mail received by the user after submission of registration info

After that the user must wait for a second e-mail which will be the final confirmation about the activation of the account. This is required in order to provide an higher security level. Once the user has received the activation confirmation, he can access the webapp by inserting e-mail address and password.

The webapp will appear as shown in Fig. 5



Data Mining & Exploration Program

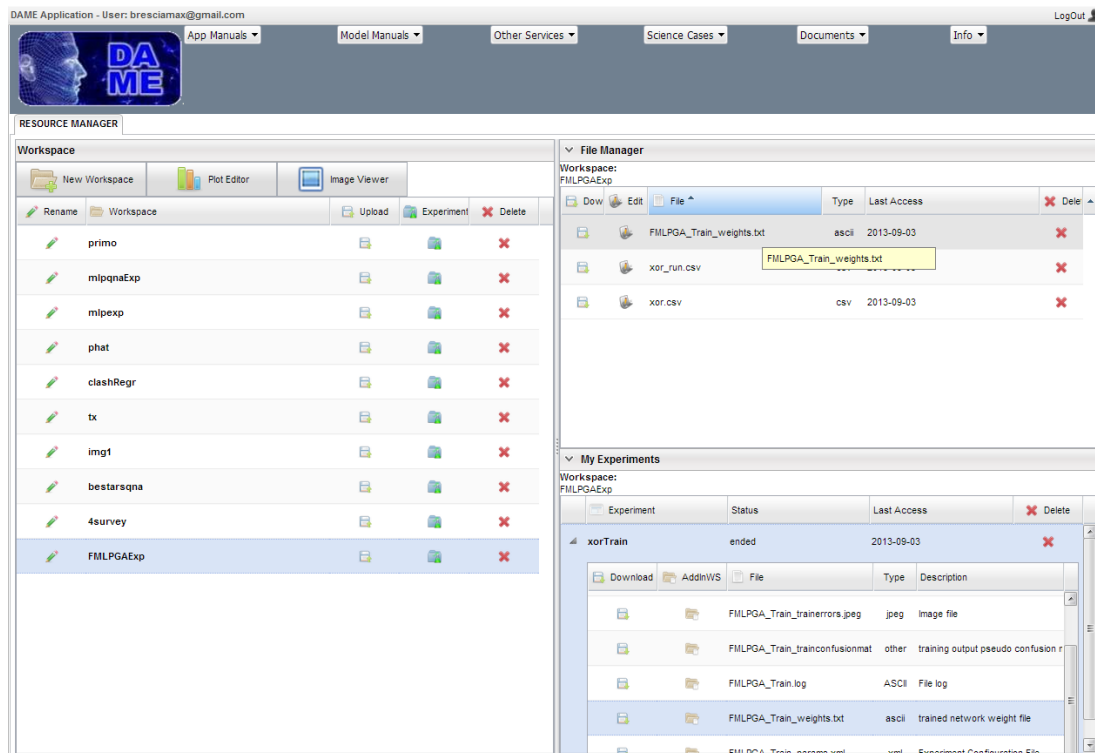


Fig. 5 – The Web Application starting main page (Resource Manager)

3.2 The command icons

The interaction between user and GUI is based on the selection of icons, which correspond to basic features available to perform actions. Here their description, related to the red circles in Fig. 6 is reported:

1. **The header menu options:** When one of the available menus is selected, a pop submenu appears with some options;
2. **Logout button:** If pressed the GUI (and related working session) is closed;
3. **Operation tabs:** The GUI is organized like a multi-tab browser. Different tabs are automatically open when user wants to edit data file to create/manipulate datasets, to upload files or to configure and launch experiments. All tabs can be closed by user, except the main one (Resource Manager);
4. **Creation of new workspaces:** When selected and named, the new workspace appears in the Workspace List Area (Workspace sub window);
5. **Workspace List Area:** portion of the main Resource Manager tab dedicated to host all user defined workspaces;
6. **Upload command:** When selected, the user is able to select a new file to be uploaded into the Workspace Data Area (Files Manager sub window). The file can be uploaded from external URI or from local (user) HD;
7. **Creation of new experiment:** When selected, the user is able to create a new experiment (a specific new tab is open to configure and launch the experiment);
8. **Rename workspace command:** When selected the user can rename the workspace;
9. **Delete Workspace command:** When selected, the user can delete the related workspace (only if no experiments are present inside, otherwise the system alerts to empty the workspace before to erase it);



DAta Mining & Exploration Program

10. **File Manager Area:** the portion of Resource Manager tab dedicated to list the data files belonging to various workspaces. All files present in this area are considered as input files for any kind of experiment;
11. **Download command:** When selected the user can download locally (on his HD) the selected file;
12. **Dataset Editor command:** When selected a new tab is open, where the user can create/edit dataset files by using all available dataset manipulation features;
13. **Delete file command:** When selected the user can delete the selected file from current workspace;
14. **Experiment List Area:** The portion of Resource Manager tab dedicated to the list of experiments and related output files present in the selected workspace;
15. **Experiment verbose list command:** When selected the user can open the experiment file list (for experiment in ended state) in a verbose mode, showing all related files created and stored;
16. **Delete Experiment command:** by clicking on it, the entire experiment (all listed files) is erased;
17. **Download experiment file command:** When selected the user can download locally (on his HD) the related experiment output file;
18. **AddinWS command:** When selected, the related file is automatically copied from the Experiment List Area to the currently active workspace File Manager Area. This feature is useful to re-use an output file of a previous experiment as input file of a new experiment (in the figure, look at the file weights.txt, that after this command is also listed in the File Manager). A file present in both areas, can be used as input either as output in the experiments;
19. **Plot Editor:** When pressed open in the resource manager four tabs: histogram, scatter 2d, scatter 3d and line plot; each tab is dedicated to a specific type of plot;
20. **Image Viewer:** When pressed open a new tab in the resource manager dedicated to the visualization of an image. The image file is intended already loaded in the File Manager.

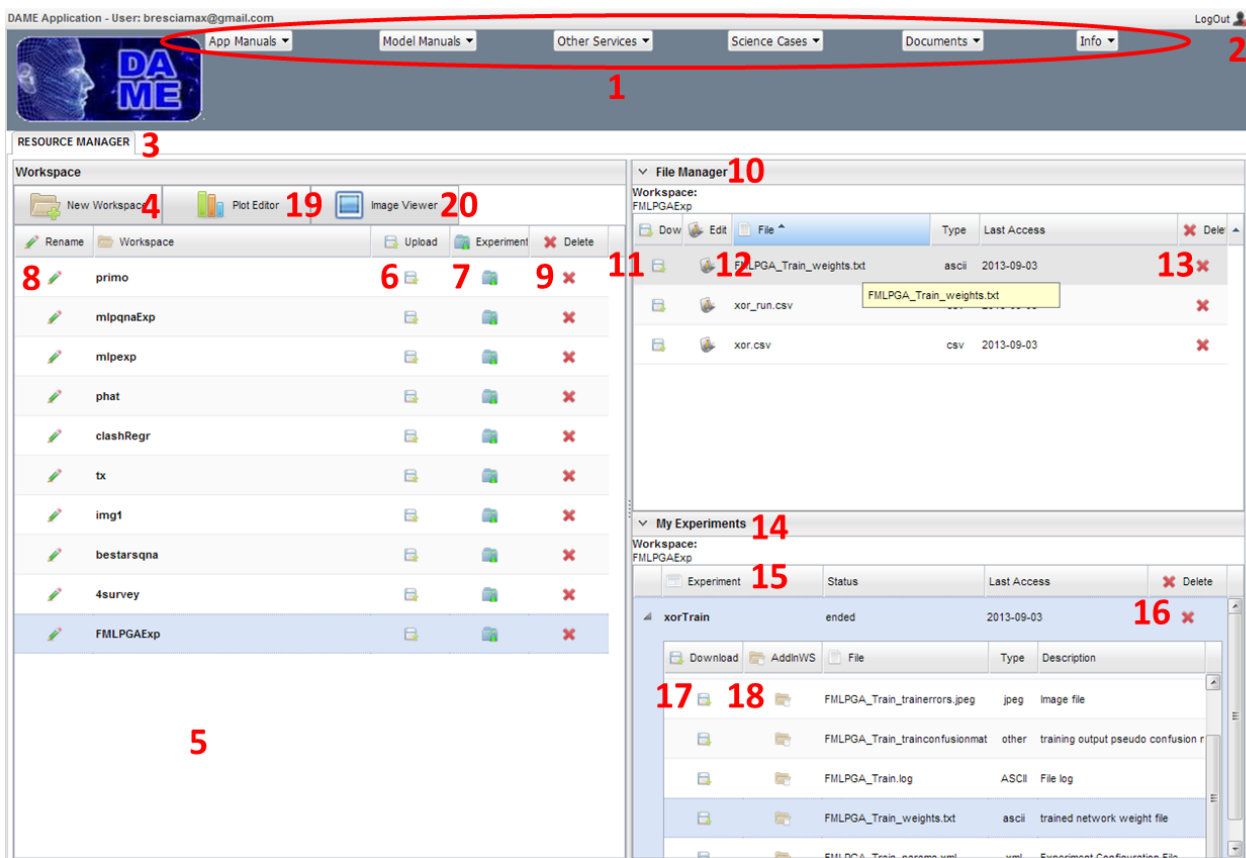


Fig. 6 – The Web Application main areas and commands



Data Mining & Exploration Program

3.3 Workspace Management

A workspace is namely a working session, in which the user can enclose resources related to scientific data mining experiments. Resources can be data files, uploaded in the workspace by the user, files resulting from some manipulations of these data files, i.e. dataset files, containing subsets of data files, selected by the user as input files for his experiments, eventually normalized or re-organized in some way (see section 3.5 for details). Resources can also be output files, i.e. obtained as results of one or more experiments configured and executed in the current “active” workspace (see section 3.7 for details).

The user can create a new or select an existing workspace, by specifying its name. After opening the workspace, this automatically becomes the “active” workspace. This means that any further action, manipulating files, configuring and executing experiments, upload/download files, will result in the active workspace, Fig. 7. In this figure it is also shown the right sequence of main actions in order to operate an experiment (workflow) in the correct way.

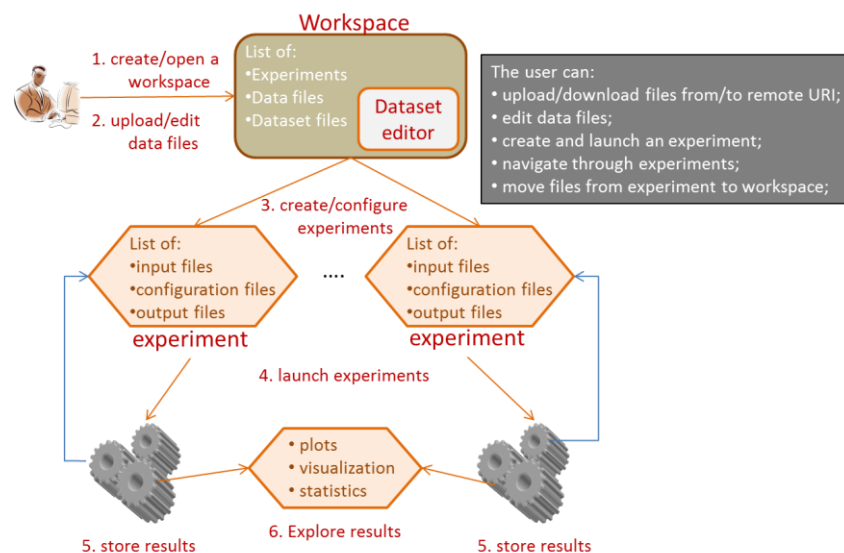


Fig. 7 – The right sequence to configure and execute an experiment workflow

So far, the basic role of a workspace is to make easier to the user the organization of experiments and related input/output files. For example the user could envelope in a same workspace all experiments related to a particular functionality domain, although using different models.

It is always possible to move (copy) files from experiment to workspace list, in order to re-use a same dataset file for multiple experiment sessions, i.e. to perform a workflow.

After access, the user must select the “active” workspace. If no workspaces are present, the user must create a new one, otherwise the user must select one of the listed workspace. The user can always create a new workspace by pressing the button as in Fig. 8.

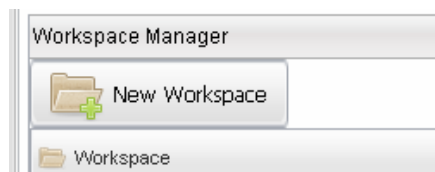


Fig. 8 – the button “New Workspace” at left corner of workspace manager window



DAta Mining & Exploration Program

As consequence the user must assign a name to the new workspace, by filling in the form field as in Fig. 9.

A small dialog box titled "New Workspace" with a close button (X) in the top right corner. It contains a text input field labeled "Name:" and two buttons at the bottom: "OK" and "Cancel".

Fig. 9 – the form field that appears after pressing the “New Workspace” button

After creation, the active workspace can be populated by data and experiments, Fig. 10.

The screenshot shows the DAME Application interface. At the top is a header bar with the DAME logo and a menu bar containing "App Manuals", "Model Manuals", "Other Services", "Science Cases", "Documents", and "Info". Below the header is the "RESOURCE MANAGER" section, which is divided into two main panels. The left panel, titled "Workspace", shows a list of workspaces: "primo", "mipqnaExp" (highlighted), "mipexp", "phat", "clashRegr", "tx", "img1", "bestarsqna", "4survey", and "FMLPGAExp". Each workspace has icons for "Rename", "Workspace", "Upload", "Experiment", and "Delete". The right panel, titled "File Manager", shows the contents of the selected workspace "mipqnaExp". It contains a table of files and a section for "My Experiments".

File	Type	Last Access	Delete
dataset_run_20.ascii	ascii	2013-06-27	X
dataset_test_20.ascii	ascii	2011-12-06	X
dataset_train_80.ascii	ascii	2011-12-06	X
dataset_train_80.ascii_mipqna_TRAIN_fr	ASCII	2013-07-02	X
mipqna_TRAIN_weights.txt	ASCII	2013-07-02	X
opt_and_struct.csv	csv	2013-07-02	X
xor.csv	csv	2011-12-05	X

Experiment	Status	Last Access	Delete
classtrain1	ended	2013-07-02	X
classtest1	ended	2013-07-02	X
classrun1	ended	2013-07-02	X
classfull1	ended	2013-07-02	X
regtrain1	ended	2013-07-02	X
regtest1	ended	2013-07-02	X
regrun1	ended	2013-07-02	X

Fig. 10 –the active workspace created in the Workspace List Area

3.4 Header Area

At the top segment of the DMS GUI there is the so-called Header Area. Apart from the DAME logo, it includes a persistent menu of options directly related to information and documentation (this document also) available online and/or addressable through specific DAME program website pages.



DAta Mining & Exploration Program

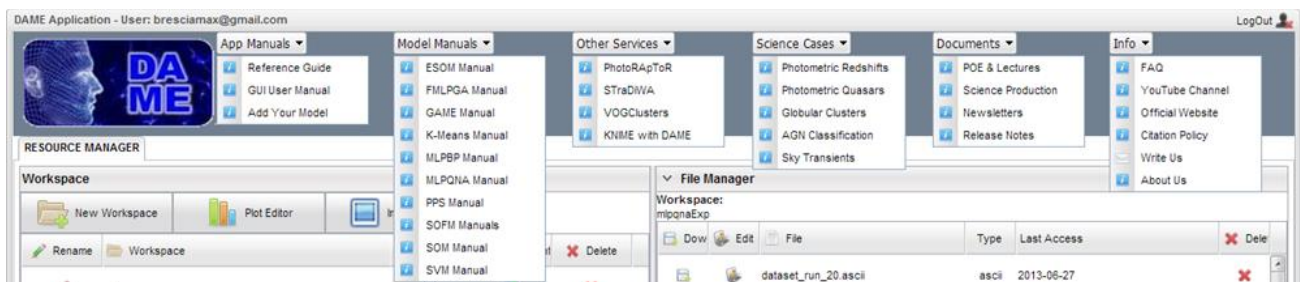


Fig. 11 –The GUI Header Area with all submenus open

The options are described in the following table (Tab. 1).

OPTIONS	HEADER	DESCRIPTION
Reference Guide	Application Manuals	http://dame.dsf.unina.it/dameware.html#appman
GUI User Manual		http://dame.dsf.unina.it/dameware.html#plugin
Extend DAME		http://dame.dsf.unina.it/dameware.html#plugin
	Model Manuals	Specific data mining model user manuals for experiments
ESOM Manual		http://dame.dsf.unina.it/dameware.html#manuals
FMLPGA Manual		http://dame.dsf.unina.it/dameware.html#manuals
MLPQNA/LEMON Manual		http://dame.dsf.unina.it/dameware.html#manuals
Random Forest Manual		http://dame.dsf.unina.it/dameware.html#manuals
K-Means Manual		http://dame.dsf.unina.it/dameware.html#manuals
MLPBP Manual		http://dame.dsf.unina.it/dameware.html#manuals
PPS Manual		http://dame.dsf.unina.it/dameware.html#manuals
SOFM Manual		http://dame.dsf.unina.it/dameware.html#manuals
SOM Manual		http://dame.dsf.unina.it/dameware.html#manuals
SVM Manual		http://dame.dsf.unina.it/dameware.html#manuals
PhotoRapToR	Other Services	http://dame.dsf.unina.it/dame_photoz.html#photoraptor
STraDIWA		http://dame.dsf.unina.it/dame_td.html
VOGCLUSTERS App		http://dame.dsf.unina.it/vogclusters.html
KNIME with DAME		http://dame.dsf.unina.it/dame_kappa.html
Photometric Redshifts	Science Cases	http://dame.dsf.unina.it/dame_photoz.html
Photometric Quasars		http://dame.dsf.unina.it/dame_qso.html
Globular Clusters		http://dame.dsf.unina.it/dame_gcs.html
AGN Classification		http://dame.dsf.unina.it/dame_agn.html
Sky Transients		http://dame.dsf.unina.it/dame_td.html
POE & Lectures	Documents	http://dame.dsf.unina.it/documents.html
Science Production		http://dame.dsf.unina.it/science_papers.html
Newsletter		http://dame.dsf.unina.it/newsletters.html
Release Notes		http://dame.dsf.unina.it/dameware.html#notes
FAQ	Info	http://dame.dsf.unina.it/dameware.html#faq
YouTube Channel		http://www.youtube.com/user/DAMEmedia
Official website		http://dame.dsf.unina.it
Citation Policy		http://dame.dsf.unina.it/#policy
Write Us		helpdame@gmail.com
About Us		http://dame.dsf.unina.it/project_members.html

Tab. 1 – Header Area Menu Options



Data Mining & Exploration Program

3.5 Data Management

The Data are the heart of the web application (data mining & exploration). All its features, directly or not, are involved within the data manipulation. So far, a special care has been devoted to features giving the opportunity to upload, download, edit, transform, submit, create data.

In the GUI input data (i.e. candidates to be inputs for scientific experiments) are basically belonging to a workspace (previously created by the user). All these data are listed in the “Files Manager” sub window. These data can be in one of the supported formats, i.e. data formats recognized by the web application as correct types that can be submitted to machine learning models to perform experiments. They are:

- **FITS (tabular and image .fits files);**
- **ASCII (.txt or .dat ordinary files);**
- **VOTable (VO compliant XML document files);**
- **CSV (Comma Separated Values .csv files);**
- **JPEG, GIF and PNG images.**

The user has to pay attention to use input data in one of these supported formats in order to launch experiments in a right way.

Other data types are permitted but not as input to experiments. For example, log, jpeg or “not supported” text files are generated as output of experiments, but only supported types can be eventually re-used as input data for experiments.

There is an exception to this rule for file format with extension **.ARFF (Attribute Relation File Format)**. These files can be uploaded and also edited by dataset editor, by using the type “CSV”. But their extension .ARFF is considered “unsupported” by the system, so you can use any of the dataset editor options to change the extension (automatically assigned as CSV). Then you can use such files as input for experiments.

These output file are generally listed in the “Experiment Manager” sub window, that can be verbosely open by the user by selecting any experiment (when it is under “ended” state).

Other data files are created by dataset creation features, a list of operations that can be performed by the user, starting from an original data file uploaded in a workspace. These data files are automatically generated with a special name as output of any of the manipulation dataset operations available.

Besides these general rules, there are some important prescriptions to take care during the preparation of data to be submitted and the setup of any Machine Learning model:

- ✓ **Input features to any machine learning model must be scalars, not arrays of values or chars. In case, you could try to find numerical representation of any not scalar or alphanumerical quantities;**
- ✓ **The input layer of a generic hierarchical neural network must be populated according to the number of physical input features of your table entries. There must be a perfect correspondence between number of input nodes and input features (columns of your table);**
- ✓ **All objects (rows) of an input table must have exactly the same number of columns. No rows with variable number of columns are allowed;**
- ✓ **Hidden layers of any multi-layer feed-forward model (i.e. layers between input and output ones) must contain a decreasing number of nodes, usually by following an empirical law: given N input nodes, the first hidden layer should have $2N+1$ nodes at least; the optional second hidden layer N-1 and so on...;**
- ✓ **For most of the available neural networks models (with exceptions of Random Forest and SVM), the class target column should be encoded with a binary representation of the class label. For**



DAta Mining & Exploration Program

example, if you have 3 different classes, you must create three different columns of targets, by encoding the 3 classes as, respectively: 001, 010, 100.

Confused? Well, don't panic please. Let's read carefully next sections.

3.5.1 Upload user data

As mentioned before, after the creation of at least one workspace, the user would like to populate the workspace with data to be submitted as input for experiments. Remember that in this section we are dealing with supported data formats only!

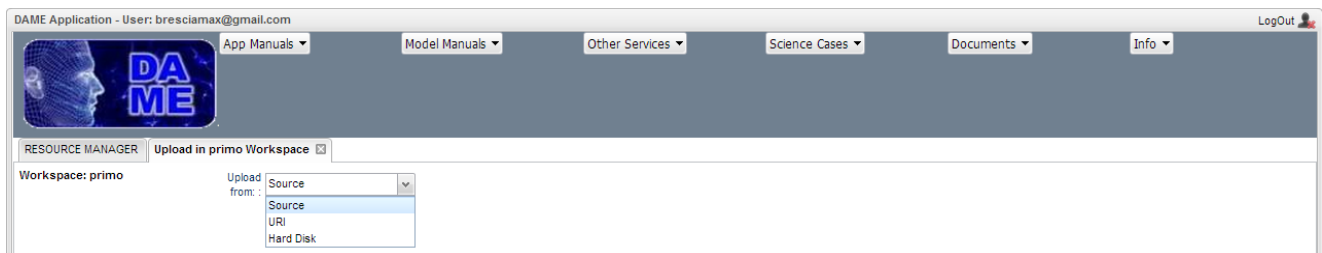


Fig. 12 – The Upload data feature open in a new tab

As shown in Fig. 12, when the user selects the “upload” command, (label nr. 6 in the Fig. 6), a new tab appears. The user can choose to upload his own data file from, respectively, from any remote URI (a priori known...!) or from his local Hard Disk.

In the first case (upload from URI), the Fig. 13 shows how to upload a supported type file from a remote address.

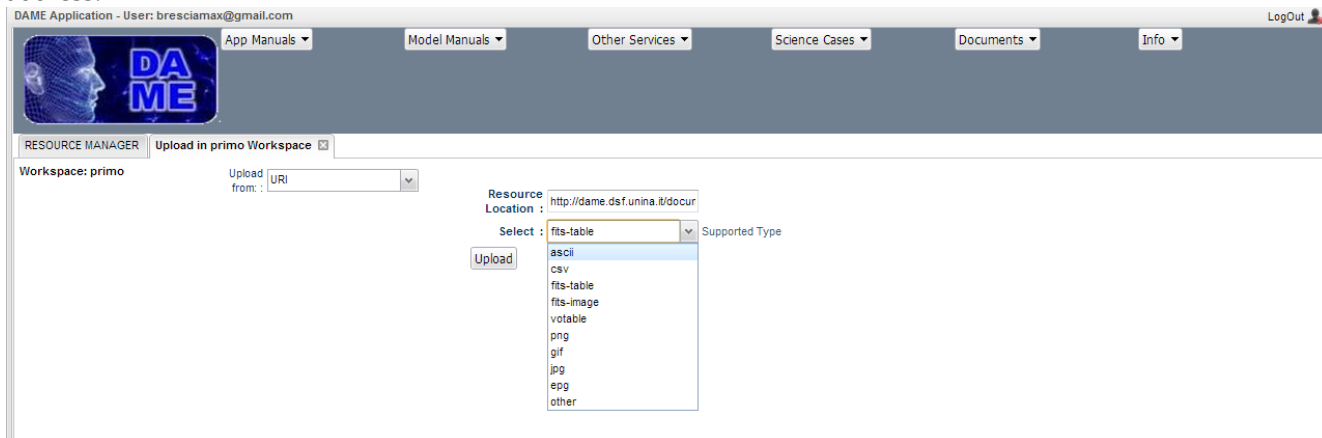


Fig. 13 – The Upload data from external URI feature

In the second case (upload from Hard Disk) the Fig. 14 shows how to select and upload any supported file in the GUI workspace from the user local HD.



Data Mining & Exploration Program

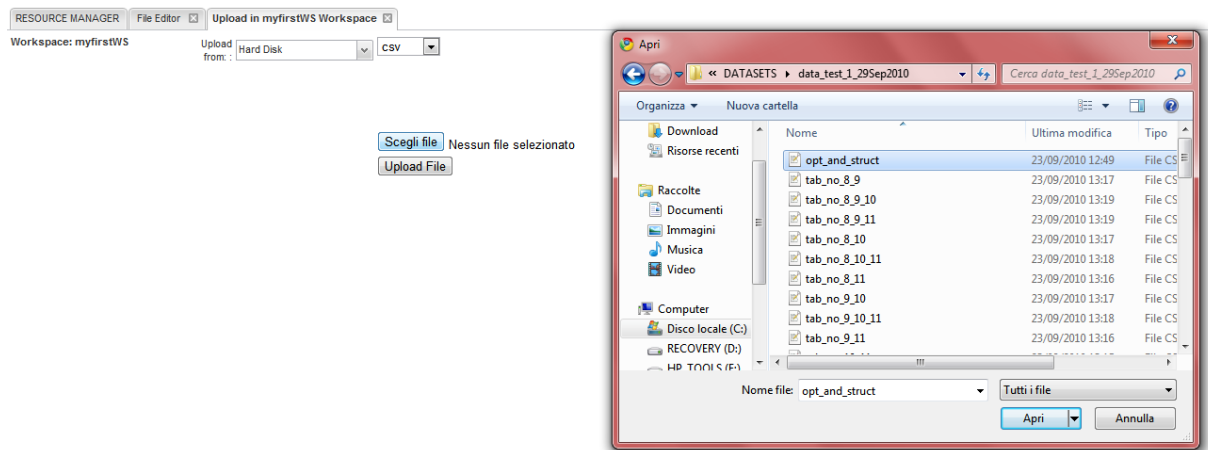


Fig. 14 – The Upload data from Hard Disk feature

After the execution of the operation, coming back to the main GUI tab, the user will find the uploaded file in the “Files Manager” sub window related with the currently active workspace, Fig. 15.

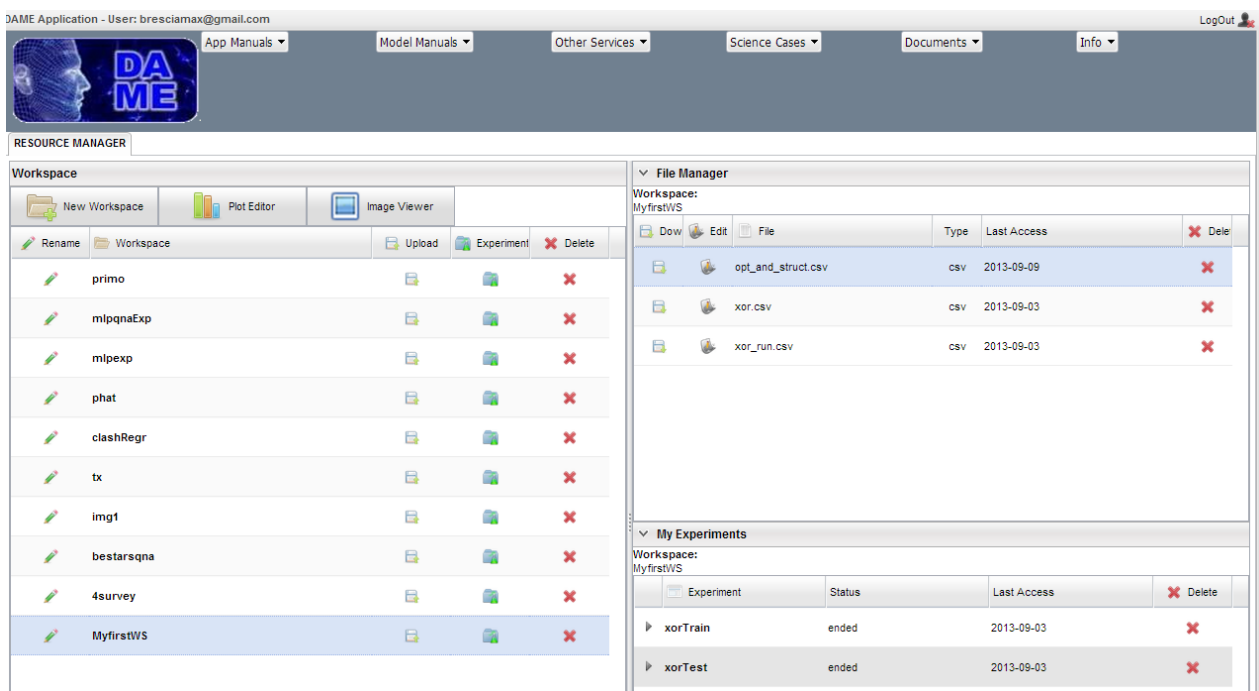


Fig. 15 – The Uploaded data file in the File Manager sub window

3.5.2 How to Create dataset files

If the user has already uploaded any supported data file in the workspace, it is possible to select it and to create datasets from it. This is a typical pre-processing phase in a machine learning based experiment, where, starting from an original data file, several different files must be prepared and provided to be submitted as input for, respectively, training, test and validate the algorithm chosen for the experiment. This pre-processing is generally made by applying one or more modification to the original data file (for example



Data Mining & Exploration Program

obtained from any astronomical observation run or cosmological simulation). The operations available in the web application are the following, Fig. 16:

- **Feature Selection;**
- **Columns Ordering;**
- **Sort Rows by Column;**
- **Column Shuffle;**
- **Row Shuffle;**
- **Split by Rows;**
- **Dataset Scale;**
- **Single Column Scale;**

All these operations, one by one, can be applied starting from a selected data file uploaded in the currently active workspace.

RESOURCE MANAGER

Upload in myfirstWS Workspace

File Editor

Workspace: myfirstWS
File: opt_and_struct_csv

Operation	Description
Feature Selection	It creates a new file, containing only the selected columns.
Columns Ordering	It creates a new file, containing the new ordering of columns, as specified in the column fields.
Sort Rows by Column	It creates a new file, containing the sorted rows, by specifying the column reference index.
Column Shuffle	It creates a new file, containing shuffled columns.
Row Shuffle	It creates a new file, containing shuffled rows.
Split by Rows	It creates as many new files as desired, each containing the specified percentage of rows. The rows distribution can be randomly extracted.
Dataset Scale	***Use only numerical Dataset*** It creates a new file, containing all data scaled in the selected range $[-1, +1]$ or $[0, +1]$.
Single Column Scale	It creates a new file, containing a single column data scaled in the selected range $[-1, +1]$ or $[0, +1]$, leaving unmodified the other columns in the original order.

Configuration

Fig. 16 – The dataset editor tab with the list of available operations

3.5.2.1 Feature Selection

This dataset operation permits to select and extract arbitrary number of columns, contained in the original data file, by saving them in a new file (of the same type and with the same extension of the original file), named as `columnSubset_<user selected name>` (i.e. with specific prefix `columnSubset`). This function is particularly useful to select training columns to be submitted to the algorithm, extracted from the whole data file. Details of the simple procedure are reported in Fig. 17 and Fig. 18.



Data Mining & Exploration Program

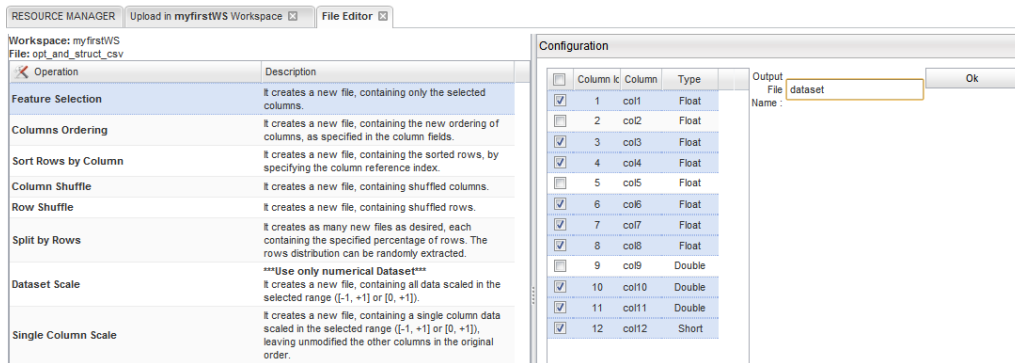


Fig. 17 – The Feature Selection operation – select columns and put saving name

As clearly visible in Fig. 17, the *Configuration* panel shows the list of columns originally present in the input data file, that can be selected by proper check boxes. Note that the whole content of the data file (in principle a massive data set) is not shown, but simply labelled by column meta-data (as originally present in the file).

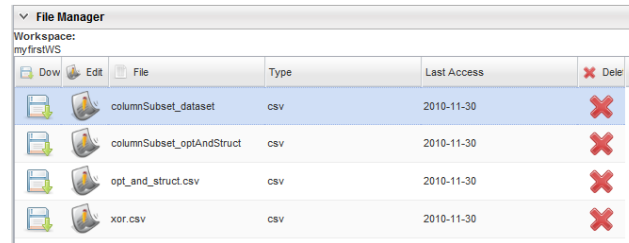


Fig. 18 –The Feature Selection operation – the new file created

3.5.2.2 Column Ordering

This dataset operation permits to select an arbitrary order of columns, contained in the original data file, by saving them in a new file (of the same type and with the same extension of the original file), named as `columnSort_<user selected name>` (i.e. with specific prefix `columnSort`). Details of the simple procedure are reported in Fig. 20.

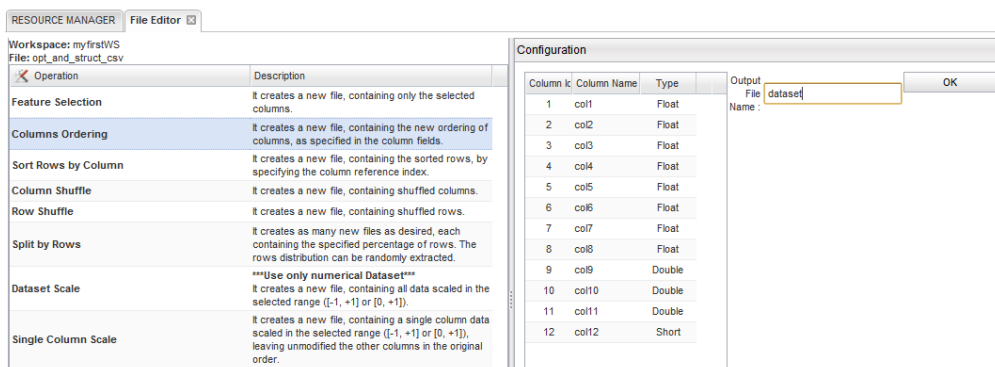


Fig. 19 –The Column Ordering operation – the starting view



Data Mining & Exploration Program

In particular, in Fig. 20 it is shown the result of several “dragging” operations operated on some columns. By selecting with mouse a column it is possible to drag it in a new desired position . At the end the new saved file will contain the new order given to data columns.

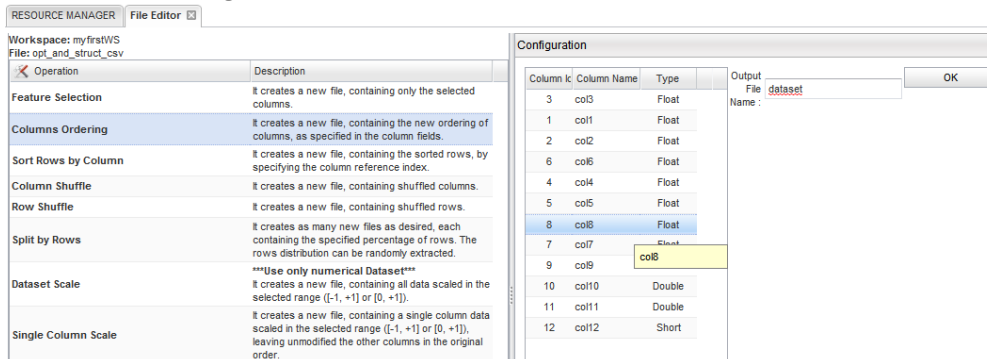


Fig. 20 –The Column Ordering operation – new order to columns

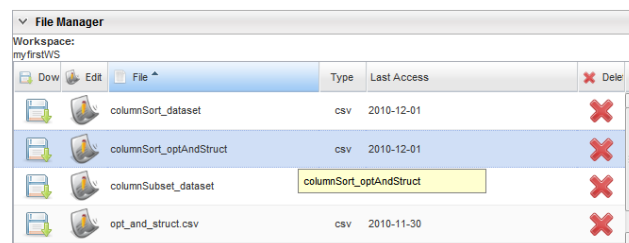


Fig. 21 – The Column Ordering operation – new file created

3.5.2.3 Sort Rows by Column

This dataset operation permits to select an arbitrary column, between those contained in the original data file, as sorting reference index for the ordering of all file rows. The result is the creation of a new file (of the same type and with the same extension of the original file), named as rowSort_<user selected name> (i.e. with specific prefix rowSort). Details of the simple procedure are reported in Fig. 22, Fig. 23 and Fig. 24.

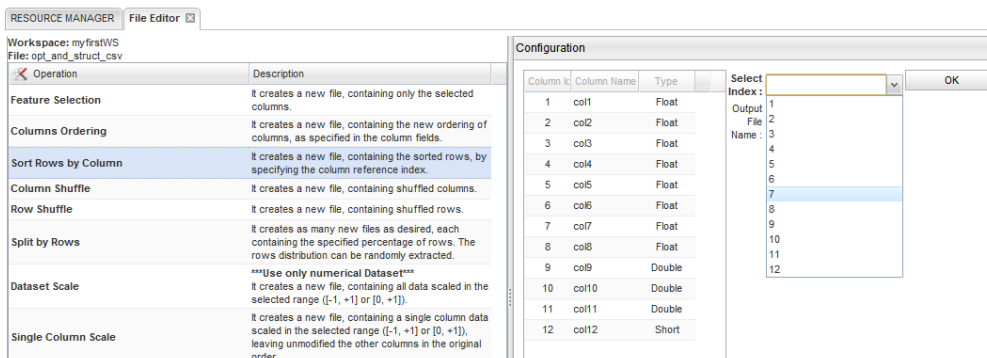


Fig. 22 –The Sort Rows by Column operation – step 1



Data Mining & Exploration Program

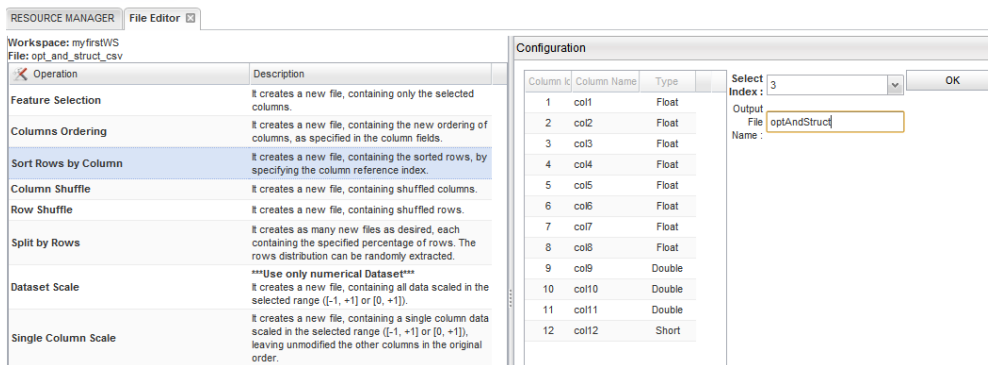


Fig. 23 –The Sort Rows by Column operation – step 2

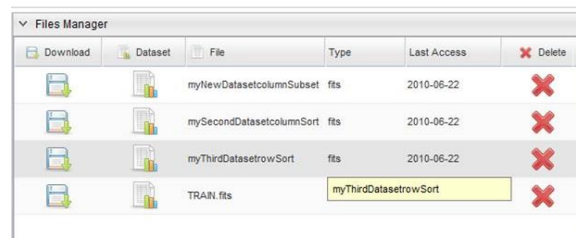


Fig. 24 –The Sort Rows by Column operation – the new file created

3.5.2.4 Column Shuffle

This dataset operation permits to operate a random shuffle of the columns, contained in the original data file. The result is the creation of a new file (of the same type and with the same extension of the original file), named as `shuffle_<user selected name>` (i.e. with specific prefix *shuffle*). Details of the simple procedure are reported in Fig. 25 and Fig. 26.

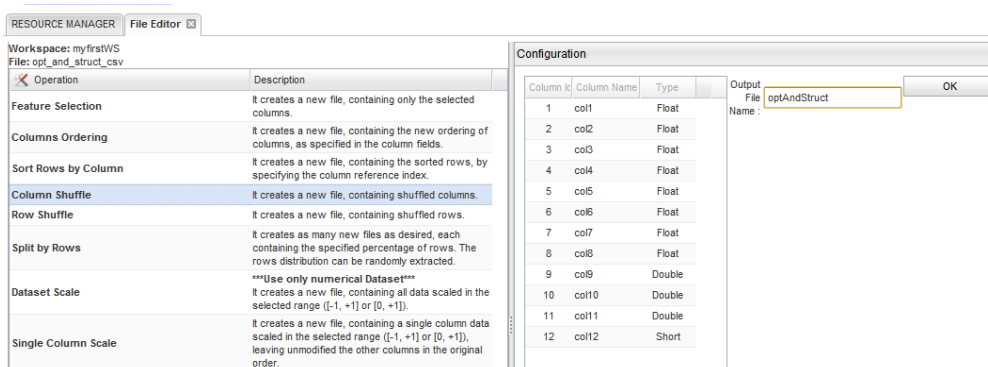


Fig. 25 –The Column Shuffle operation – step 1



Data Mining & Exploration Program

File Manager

Workspace: myfirstWS

Down	Up	Edit	File	Type	Last Access	Del
			columnSort_optAndStruct	csv	2010-12-01	
			columnSubset_dataset	csv	2010-11-30	
			opt_and_struct.csv	csv	2010-11-30	
			rowShuffle_optAndStruct	csv	2010-12-01	
			rowSort_optAndStruct	csv	2010-12-01	
			shuffle_optAndStruct	csv	2010-12-01	
			xor.csv	csv	2010-11-30	

Fig. 26 –The Column Shuffle operation – the new file created

3.5.2.5 Row Shuffle

This dataset operation permits to operate a random shuffle of the rows, contained in the original data file. The result is the creation of a new file (of the same type and with the same extension of the original file), named as rowShuffle_<user selected name> (i.e. with specific prefix *rowShuffle*). Details of the simple procedure are reported in Fig. 27 and Fig. 28.

RESOURCE MANAGER

File Editor

Workspace: myfirstWS

File: opt_and_struct.csv

Operation	Description
Feature Selection	It creates a new file, containing only the selected columns.
Columns Ordering	It creates a new file, containing the new ordering of columns, as specified in the column fields.
Sort Rows by Column	It creates a new file, containing the sorted rows, by specifying the column reference index.
Column Shuffle	It creates a new file, containing shuffled columns.
Row Shuffle	It creates a new file, containing shuffled rows.
Split by Rows	It creates as many new files as desired, each containing the specified percentage of rows. The rows distribution can be randomly extracted.
Dataset Scale	***Use only numerical Dataset*** It creates a new file, containing all data scaled in the selected range [-1, +1] or [0, +1].
Single Column Scale	It creates a new file, containing a single column data scaled in the selected range [-1, +1] or [0, +1], leaving unmodified the other columns in the original order.

Configuration

Column Id	Column Name	Type
1	col1	Float
2	col2	Float
3	col3	Float
4	col4	Float
5	col5	Float
6	col6	Float
7	col7	Float
8	col8	Float
9	col9	Double
10	col10	Double
11	col11	Double
12	col12	Short

Output File Name:

optAndStruct

OK

Fig. 27 –The Row Shuffle operation – step 1

File Manager						
Workspace: myfirstWS						
Down	Up	Edit	File	Type	Last Access	Del
			columnSort_dataset	csv	2010-12-01	
			columnSort_optAndStruct	csv	2010-12-01	
			columnSubset_dataset	csv	2010-11-30	
			opt_and_struct.csv	csv	2010-11-30	
			rowShuffle_optAndStruct	csv	2010-12-01	
			rowSort_optAndStruct	csv	2010-12-01	
			shuffle_optAndStruct	csv	2010-12-01	

Fig. 28 –The Row Shuffle operation – the new file created



DAta Mining & Exploration Program

3.5.2.6 Split by Rows

This dataset operation permits to split the original file into two new files containing the selected percentages of rows, as indicated by the user. The user can move one of the two sliding bars in order to fix the desired percentage. The other sliding bar will automatically move in the right percentage position. The new file names are those filled in by the user in the proper name fields as *split1_<user selected name>*(*split2_<user selected name>*) (i.e. with specific prefix *split1* and *split2*). Details of the simple procedure are reported in Fig. 29, Fig. 30.

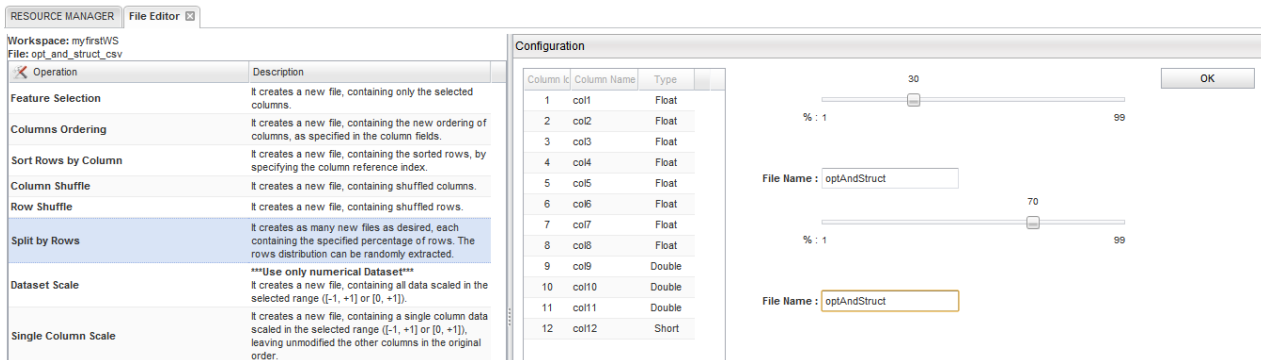


Fig. 29 –The Split by Rows operation – step 1

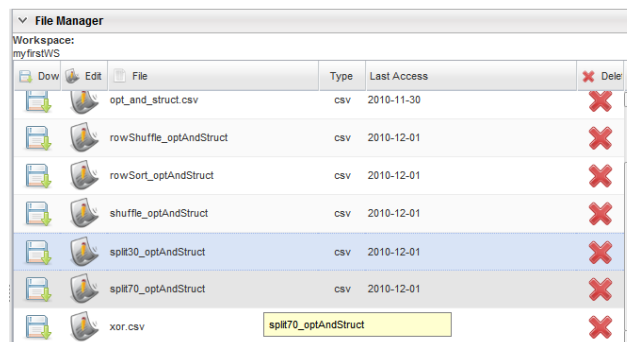


Fig. 30 –The Split by Rows operation – the new files created

3.5.2.7 Dataset Scale

This dataset operation (that works on numerical data files only!) permits to normalize column data in one of two possible ranges, respectively, $[-1, +1]$ or $[0, +1]$. This is particularly frequent in machine learning experiments to submit normalized data, in order to achieve a correct training of internal patterns. The result is the creation of a new file (of the same type and with the same extension of the original file), named as *scale_<user selected name>* (i.e. with specific prefix *scale*). Details are reported in Fig. 31.



Data Mining & Exploration Program

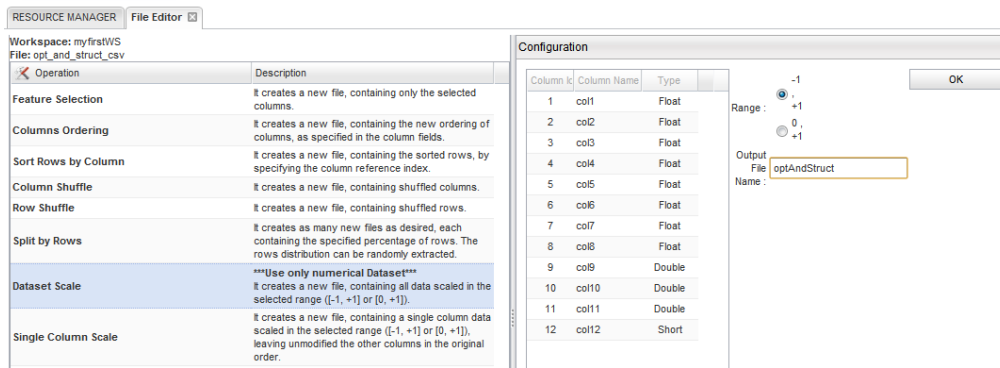


Fig. 31 –The Dataset Scale operation – step 1

3.5.2.8 Single Column Scale

This dataset operation (that works on numerical data files only!) permits to normalize a single selected column, between those contained in the original file, in one of two possible ranges, respectively, $[-1, +1]$ or $[0, +1]$. The result is the creation of a new file (of the same type and with the same extension of the original file), named as `scaleOneCol_<user selected name>` (i.e. with specific prefix `scaleOneCol`). Details of the simple procedure are reported in Fig. 32.

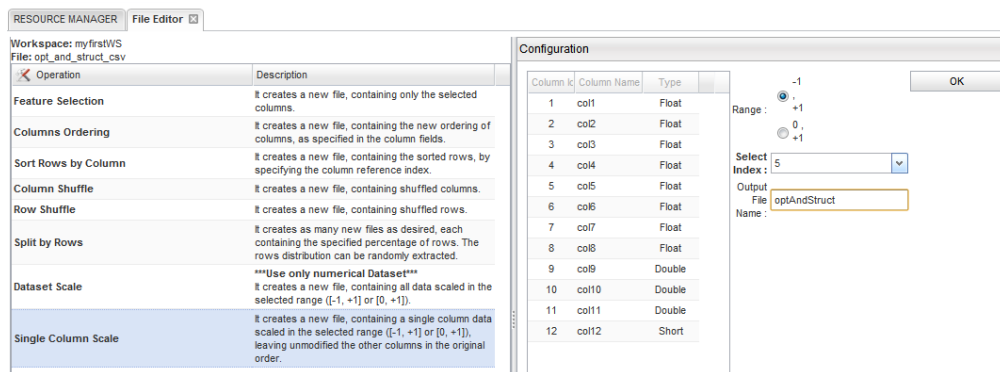


Fig. 32 –The Single Column Scale operation – step 1

3.5.3 Download data

All data files (not only those of supported type) listed in the workspace and/or in the experiment panels, respectively, “Files Manager” and “Experiment Manager”, can be downloaded by the user on his own hard disk, by simply selecting the icon labelled with “Download” in the mentioned panels.

3.5.4 Moving data files

The virtual separation of user data files between workspace and experiment files, located in the respective panels (“File Manager” for workspace files, and “My Experiments” for experiment files), is due to the



DAta Mining & Exploration Program

different origin of such files and depends on their registration policy into the web application database. The data files present in the workspace list (“File Manager” area panel) are usually registered as “input” files, i.e. to be submitted as inputs for experiments. While others, present in the experiment list (“My Experiments” panel), are considered as “output” files, i.e. generated by the web application after the execution of an experiment.

It is not rare, in machine learning complex workflows, to re-use some output files, obtained after training phase, as inputs of a test/validation phase of the same workflow. This is true for example for a MLP weight matrix file, output of the training phase, to be re-used as input weight matrix of a test (or validation) session of the same network.

In order to make available this fundamental feature in our application, the icon command nr. 18 (AddInWS) in Fig. 6, associated to each output file of an experiment, can be selected by the user in order to “copy” the file from experiment output list to the workspace input list, becoming immediately available as input file for new experiments belonging to the same workspace: **as important remark, in the beta release it is not yet possible to “move” files from a workspace to another.** The alternative procedure to perform this action is to download the file on user local Hard Disk and to upload it into another desired workspace in the webapp.

3.6 Plotting and Visualization

The final release of the web application offers two new options for plotting and visualization of data (tables or images).

3.6.1 Plotting

By pressing the “Plot Editor” button in the main menu a series of plot tabs will appear. Each one is dedicated to a specific type of plot, for instance Histogram, Scatter Plot 2D, Scatter Plot 3D and Line Plot.

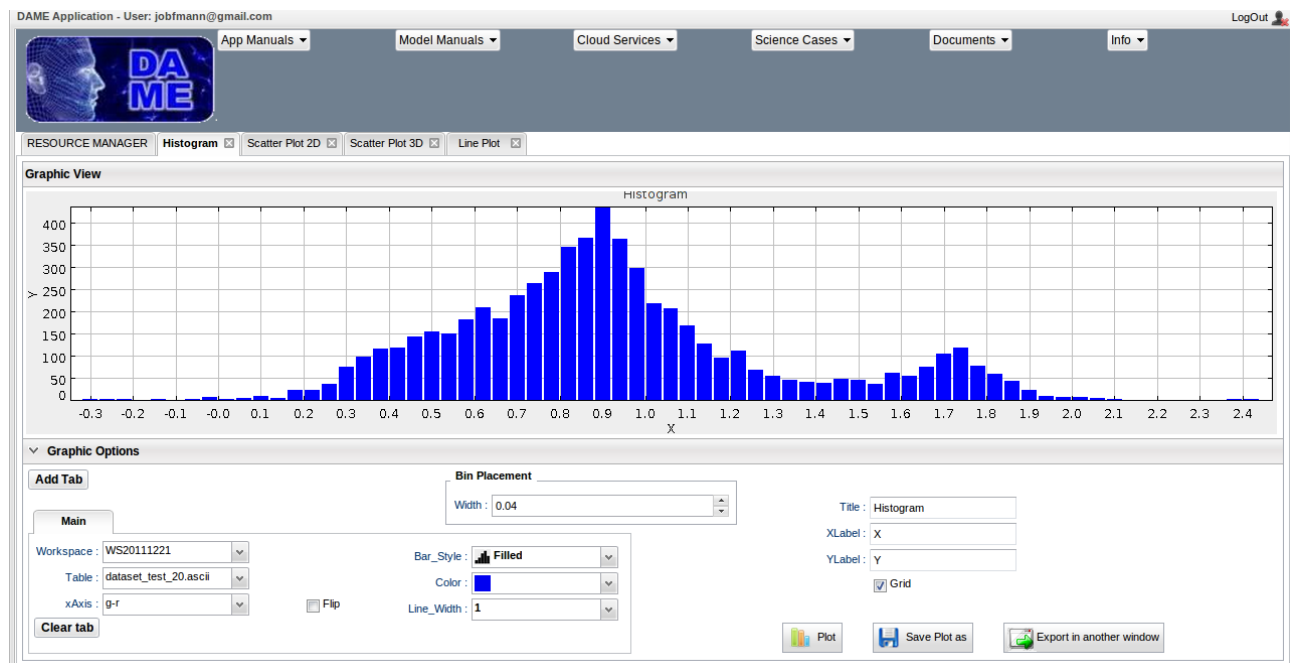


Fig. 33 – The Histogram tab



DAta Mining & Exploration Program

As shown in Fig. 33 there is the possibility to create and visualize an histogram of any table file previously loaded or produced in the web application.

There are several options:

- Workspace: the user workspace hosting the table;
- Table: the name of the table to be plotted;
- xAxis: selection of the column of table to be plotted;
- Bar_Style: style of the bars;
- Color: color of the plot bars;
- Line_Width: width of the bars;
- Flip: enable the flipping of the x Axis of the plot;
- Title: title of the plot;
- Xlabel: label of the x axis;
- Ylabel: label of the y axis;
- Grid: enable/disable the grid in the plot;
- Bin Placement: change the bin width of plot;
- Clear Tab: clear the tab;
- Plot: creation and visualization of the selected histogram;
- Save Plot As: plot saving with user typed name;
- Export in another window: the plot will be moved in an independent tab of the web browser;
- Add Tab: enable the creation of a multi layer histogram as shown in Fig. 34.

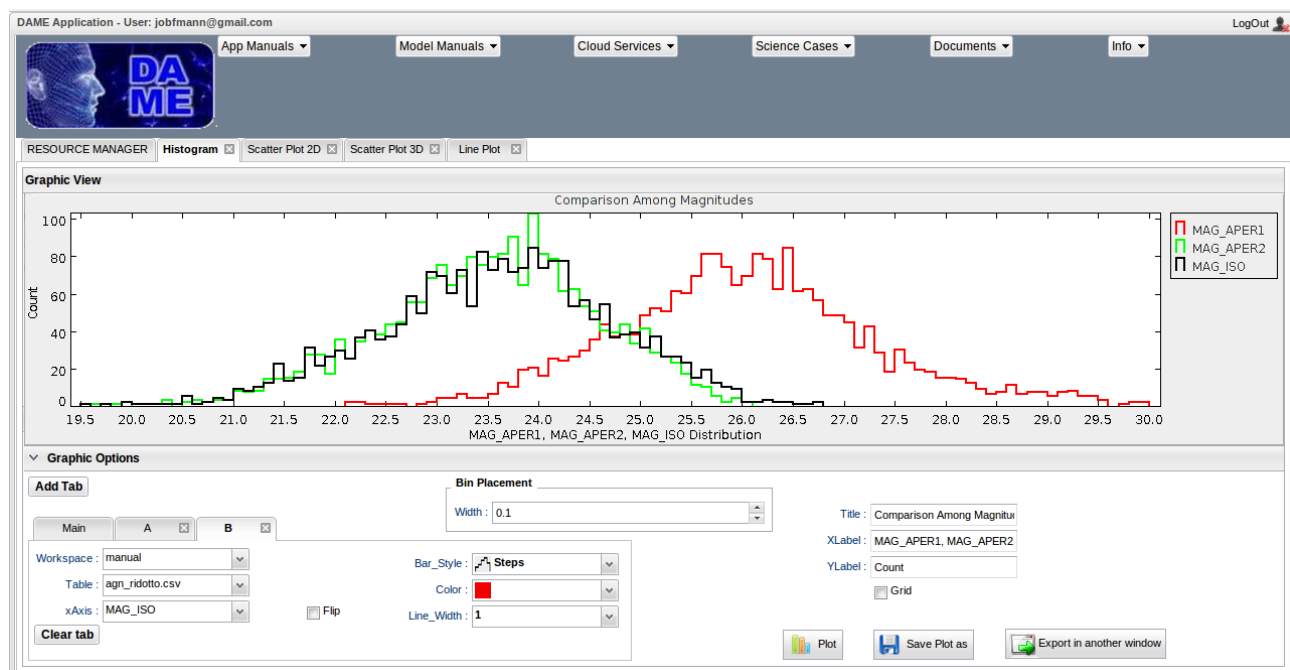


Fig. 34 – A multi layer histogram plot

ADVERTISEMENT: whenever the user change any parameter of the current plot, it is needed to click the button “Plot” to refresh the visualized plot.



DAta Mining & Exploration Program

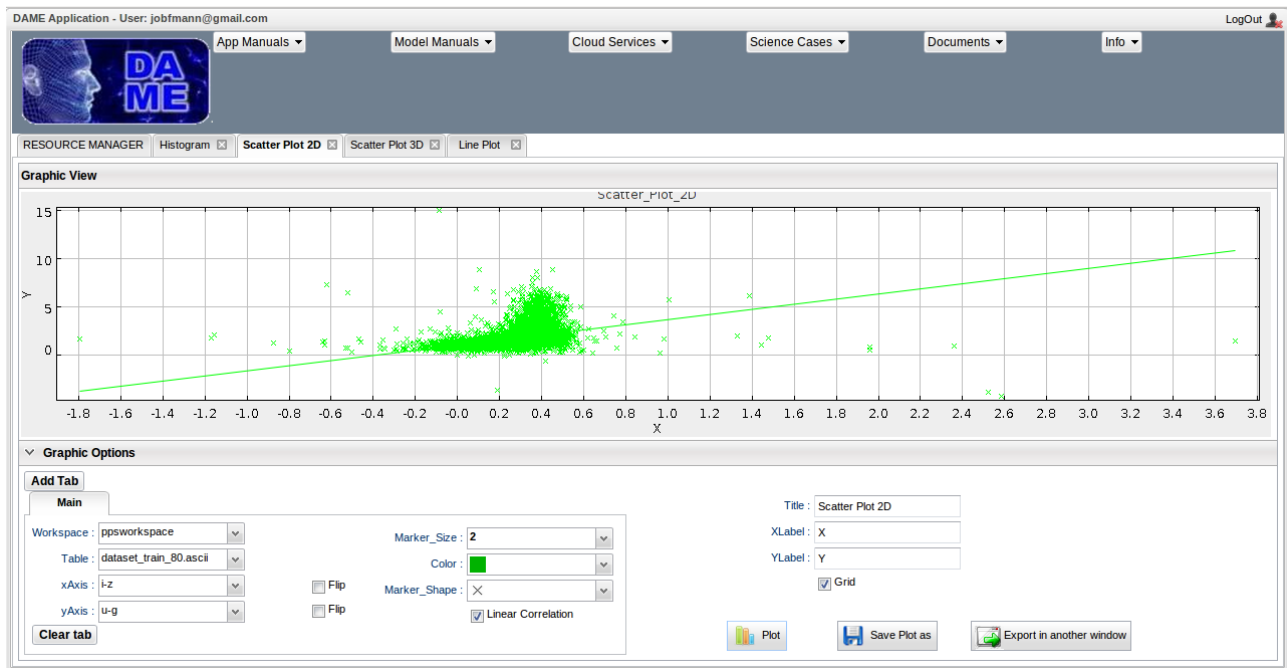


Fig. 35 – The Scatter 2D tab

As shown in Fig. 35 there is the possibility to create and visualize a scatter 2D of any table file previously loaded or produced in the web application.

There are several options:

- Workspace: the user workspace hosting the table;
- Table: the name of the table to be plotted;
- xAxis: selection of the x column of table to be plotted;
- yAxis: selection of the y column of table to be plotted;
- Marker_Size: size of the marker;
- Color: color of the plot bars;
- Marker_Shape: shape of the marker;
- Line_Width: width of the bars;
- Linear Correlation: enable the drawing of a line based on linear correlation of columns;
- Flip: enable the flipping of the x Axis of the plot;
- Flip: enable the flipping of the y Axis of the plot;
- Title: title of the plot;
- Xlabel: label of the x axis;
- Ylabel: label of the y axis;
- Grid: enable/disable the grid in the plot;
- Clear Tab: clear the tab;
- Plot: creation and visualization of the selected histogram;
- Save Plot As: plot saving with user typed name;
- Export in another window: the plot will be moved in an independent tab of the web browser;
- Add Tab: enable the creation of a multi layer scatter 2D as shown in Fig. 36.



Data Mining & Exploration Program

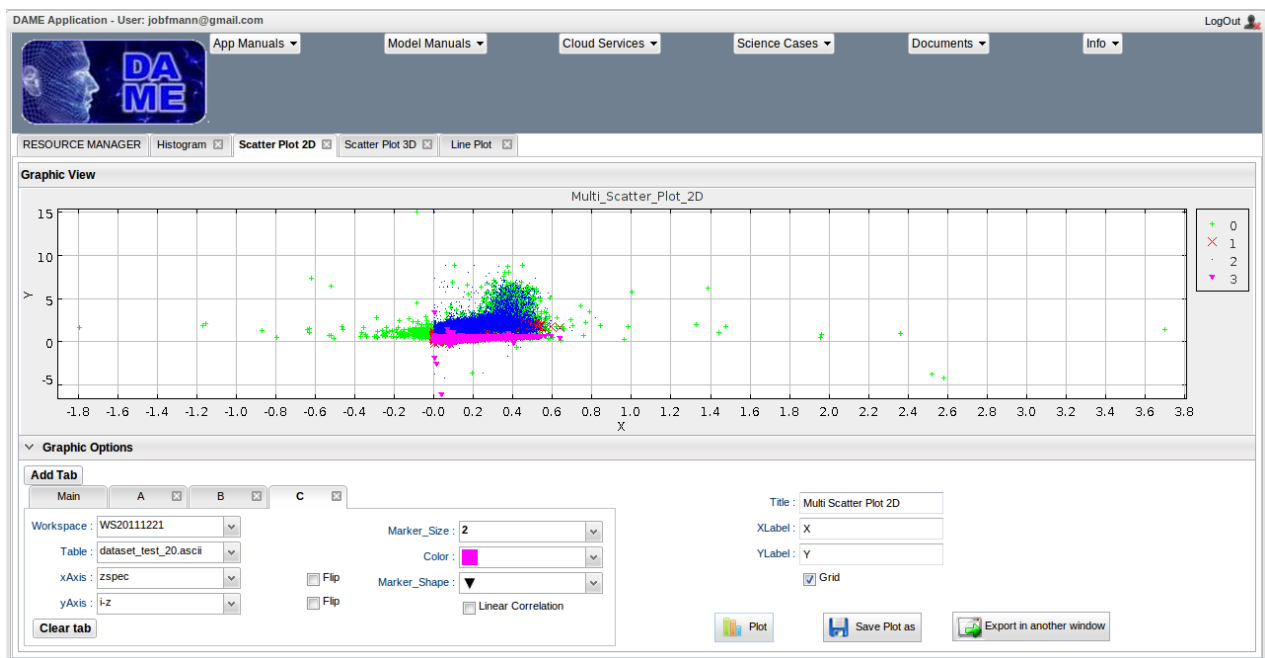


Fig. 36 – A multi layer scatter 2D plot

ADVERTISEMENT: whenever the user change any parameter of the current plot, it is needed to click the button “Plot” to refresh the visualized plot.

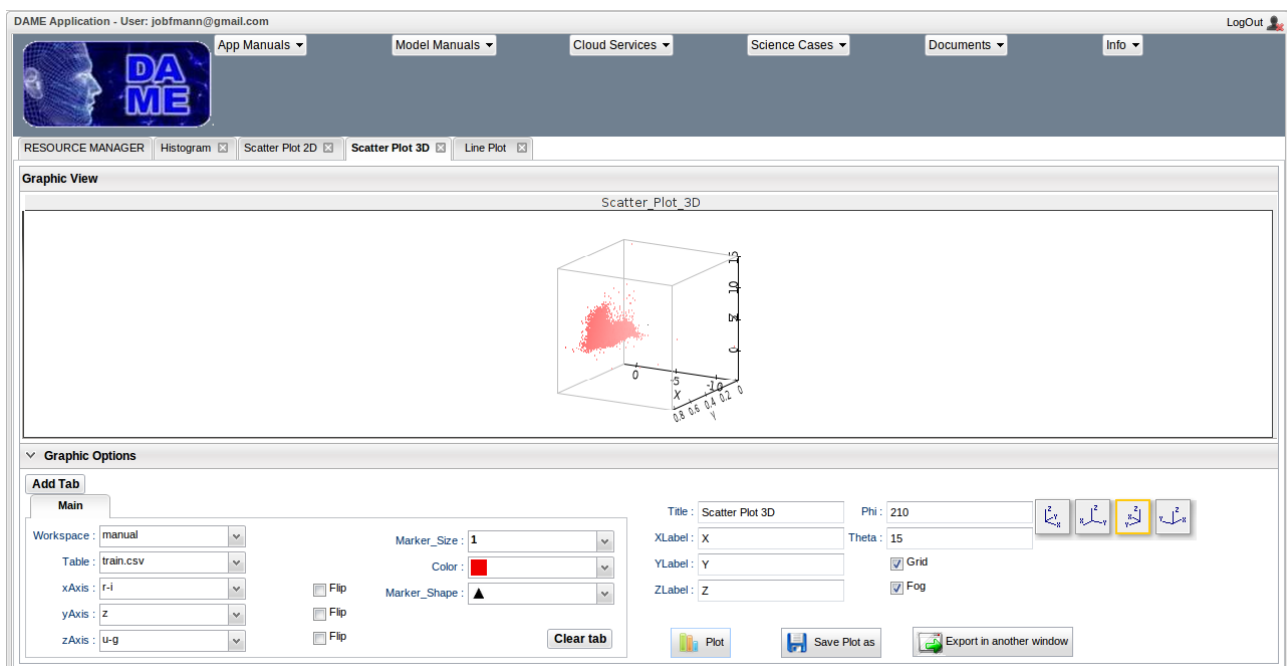


Fig. 37 – The Scatter Plot 3D tab

As shown in Fig. 37 there is the possibility to create and visualize a scatter 3D of any table file previously loaded or produced in the web application.

There are several options:

- Workspace: the user workspace hosting the table;
- Table: the name of the table to be plotted;



DAta Mining & Exploration Program

- xAxis: selection of the x column of table to be plotted;
- yAxis: selection of the y column of table to be plotted;
- zAxis: selection of the z column of table to be plotted;
- Marker_Size: size of the marker;
- Color: color of the plot bars;
- Marker_Shape: shape of the marker;
- Line_Width: width of the bars;
- Flip: enable the flipping of the x Axis of the plot;
- Flip: enable the flipping of the y Axis of the plot;
- Flip: enable the flipping of the z Axis of the plot;
- Title: title of the plot;
- Xlabel: label of the x axis;
- Ylabel: label of the y axis;
- Zlabel: label of the z axis;
- Grid: enable/disable the grid in the plot;
- Fog: enable/disable the fog effect;
- Phi: rotation angle in degrees;
- Theta: rotation angle in degrees;
- Orientation buttons: four predefined couples of Phi and Theta;
- Clear Tab: clear the tab;
- Plot: creation and visualization of the selected histogram;
- Save Plot As: plot saving with user typed name;
- Export in another window: the plot will be moved in an independent tab of the web browser;
- Add Tab: enable the creation of a multi layer scatter plot 3D.

ADVERTISEMENT: whenever the user change any parameter of the current plot, it is needed to click the button “Plot” to refresh the visualized plot.

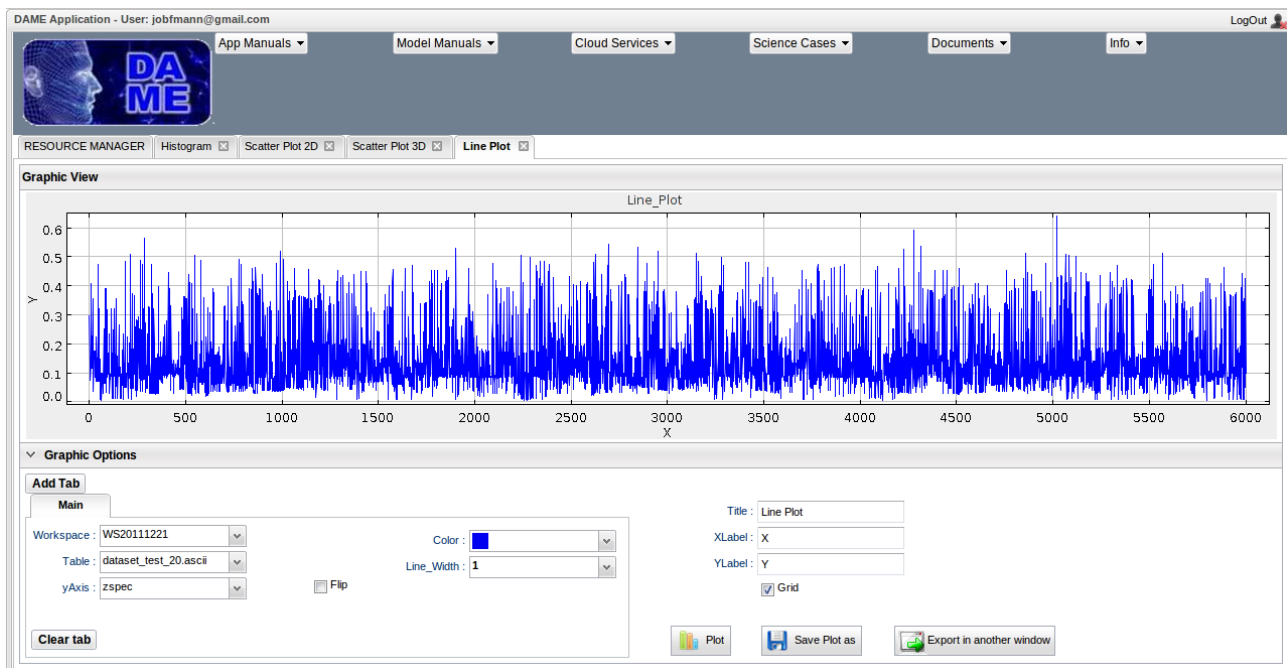


Fig. 38 – The Line Plot tab



DAta Mining & Exploration Program

As shown in Fig. 38 there is the possibility to create and visualize an histogram of any table file previously loaded or produced in the web application.

There are several options:

- Workspace: the user workspace hosting the table;
- Table: the name of the table to be plotted;
- yAxis: selection of the column of table to be plotted;
- Color: color of the plot line;
- Line_Width: width of the line;
- Flip: enable the flipping of the x Axis of the plot;
- Title: title of the plot;
- Xlabel: label of the x axis;
- Ylabel: label of the y axis;
- Grid: enable/disable the grid in the plot;
- Clear Tab: clear the tab;
- Plot: creation and visualization of the selected histogram;
- Save Plot As: plot saving with user typed name;
- Export in another window: the plot will be moved in an independent tab of the web browser;
- Add Tab: enable the creation of a multi layer line plot as shown in Fig. 39.

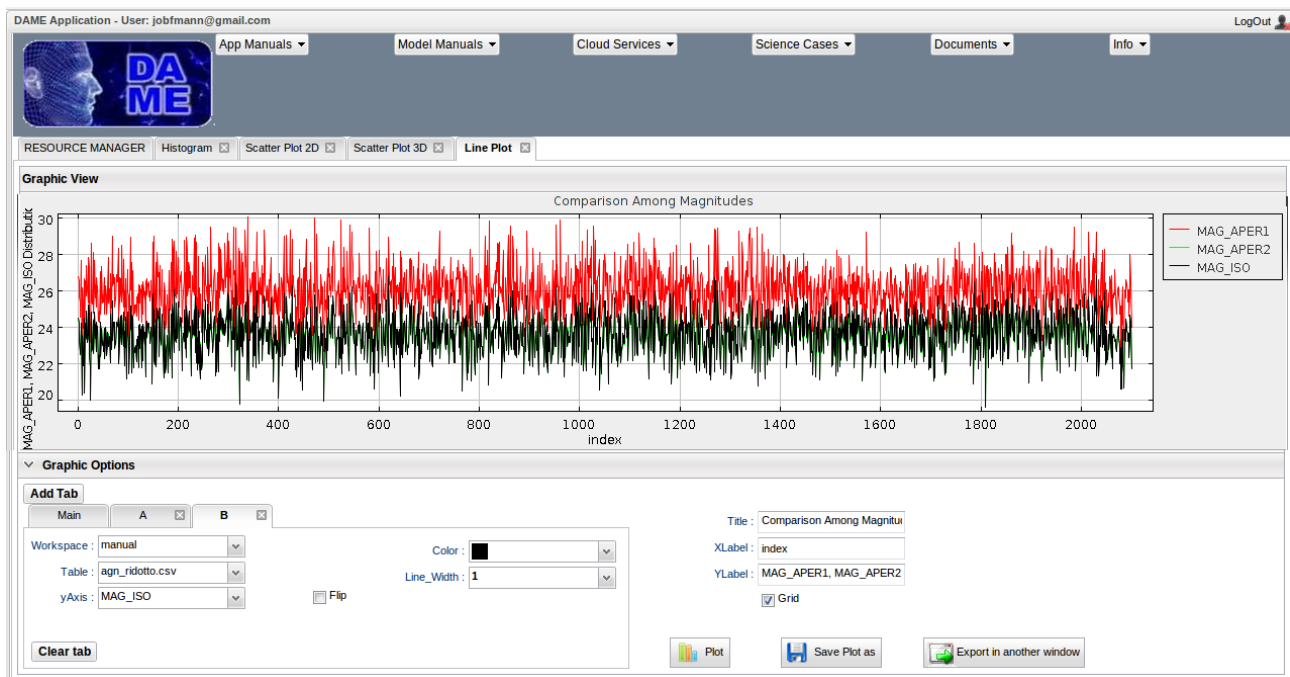


Fig. 39 – A multi layer line plot

ADVERTISEMENT: whenever the user change any parameter of the current plot, it is needed to click the button “Plot” to refresh the visualized plot.



Data Mining & Exploration Program

3.6.2 Visualization

This option can be enabled from the main tab of the GUI by simply clicking on the menu button “Image Viewer”. A dedicated tab will appear in the Resource Manager giving the possibility to load and visualize any image previously uploaded or produced in the web application.

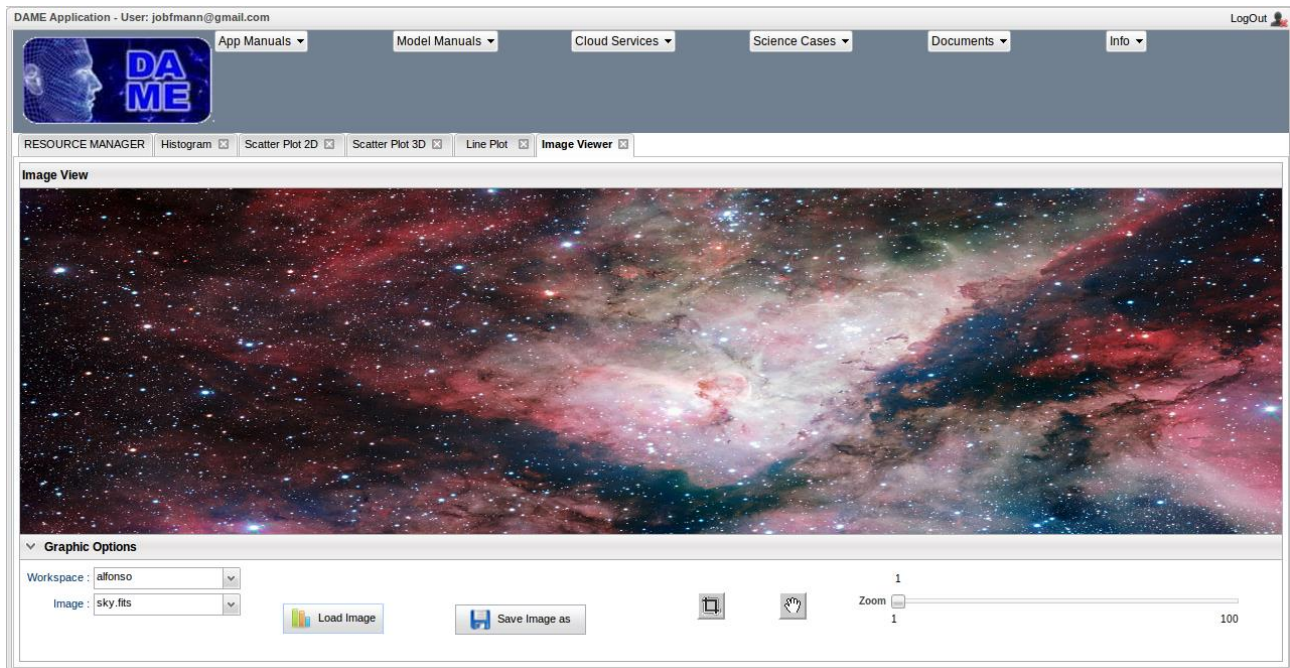


Fig. 40 – The visualization tab

As shown in Fig. 40 the visualization tab offer the following options:

- Workspace: the user workspace hosting the image;
- Image: the name of the image to be visualized;
- Load Image: after the selection of workspace and image this button shows the image;
- Crop: button used to crop the image;
- Hand: button used to move the image;
- Zoom: sliding bar used to zoom the image;
- Save Image as: button used to save the modified image.

ADVERTISEMENT: multi image fits files are not supported by this functionality.

3.7 Experiment Management

After creating at least one workspace, populating it with input data files (of supported type) and optionally creating any dataset file, the next logical operation required is the configuration and launch of an experiment. In what follows, we will explain the experiment configuration and execution by making use of an example (very simple not linearly separable XOR problem) which can be replicated by the user by using the xor.csv and xor_run.csv data files (downloadable from the beta intro web page, <http://dame.dsf.unina.it/dameware.html>).



Data Mining & Exploration Program

The Fig. 41 shows the initial step required, i.e. the selection of the icon command nr. 7 of Fig. 6 in order to create the new experiment.

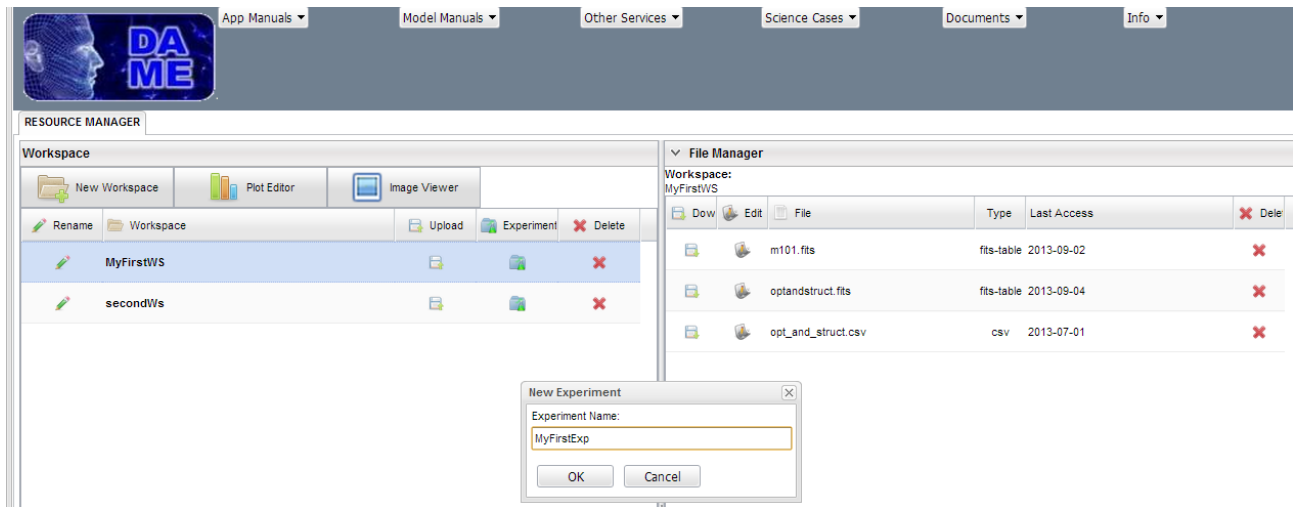


Fig. 41 – Creating a new experiment (by selecting icon “Experiment” in the workspace)

Immediately after, an automatic new tab appears, making available all basic features to select, configure and launch the experiment. In particular there is the list of couples [functionality]-[model] to choose for the current experiment. The proper choice should be done in order to solve a particular problem. It depends basically on the dataset to be used as input and on the output the user wants to obtain. Please, refer to the particular model reference manual for more details.

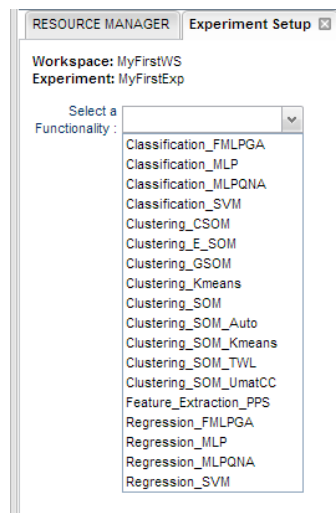


Fig. 42 – The new tab open after creation of a new experiment with the list of available options

The user can choose between classification, regression or clustering type of functionality to be applied to his problem. Each of these functionalities can be achieved by associating a particular data mining model, chosen between following types:

- **MLP** : Multilayer Perceptron neural network trained by standard Back Propagation (descent gradient of the error) learning rule. Associated functionalities are classification and regression;



Data Mining & Exploration Program

- **FMLPGA:** Fast Multilayer Perceptron neural network trained by Genetic Algorithm learning rule. Associated functionalities are classification and regression. This model is available in two versions: CPU and GPU; the second one is the parallelized version;
- **SVM:** Support Vector Machine model. Associated functionalities are classification and regression;
- **MLPQNA:** Multilayer Perceptron neural network trained by Quasi Newton learning rule. Associated functionalities are classification and regression;
- **LEMON:** Multilayer Perceptron neural network trained by Levenberg-Marquardt learning rule. Associated functionalities are classification and regression;
- **RANDOM FOREST:** Randomly generated forest of decision trees network. Associated functionalities are classification and regression;
- **KMEANS:** Standard Kmeans algorithm. Associated functionality is clustering;
- **CSOM:** Customized Self Organizing Feature Map for clustering on FITS images;
- **GSOM:** Gated Self Organizing Map (SOFM) for clustering on text and/or image files;
- **PPS:** Probabilistic Principal Surfaces for feature extraction;
- **SOM:** Self organizing Map for pre-clustering on text or image files;
- **SOM + Auto:** SOM with an automatized post processing phase for clustering on text or image files;
- **SOM + Kmeans:** SOM with a Kmeans based post processing phase for clustering on text or image files;
- **SOM + TWL:** SOM with a Two Winners Linkage (TWL) based post processing phase for clustering on text or image files;
- **SOM + UmatCC:** SOM with an U-matrix Connected Components (UmatCC) based post processing phase for clustering on text or image files;
- **ESOM:** Evolving SOM for pre-clustering on text or image files.

Specific related manuals are available to obtain detailed information about the use of the above models (see webapp header menu options).

After the selection of the proper functionality-model, the tab will show (greyed) some options and the possibility to select the use case. The greyed options (like help button) will be activated after the selection of the use case to be configured and launched.

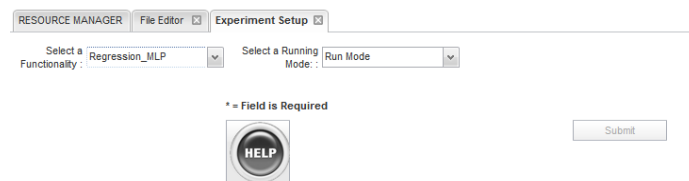


Fig. 43 – The new state of the experiment configuration tab after the selection of the model

As known, data mining models, following machine learning paradigm, offer a series of use cases (see figures below):

- **Train:** training (learning) phase in which the model is trained with the user available BoK;



DAta Mining & Exploration Program

Fig. 44 – The configuration options in the Train use case

- **Test**: a sort of validation of the training phase. It can done by submitting the same training dataset, or a subset or a mix between already submitted and new dataset patterns;

Fig. 45 – The configuration options in the Test use case

- **Run**: normal use of the already trained model;

Fig. 46 – The configuration options in the Run use case

- **Full**: the complete and automatic serialized execution of the three previous use cases (train, test and Run). It is a sort of workflow, considered as a complete and exhaustive experiment for a specific problem.



Data Mining & Exploration Program

Fig. 47 – The configuration options in the Full use case

In all the above use case tabs, the help button redirects to a specific web page, reporting in verbose mode detailed description of all parameters. In particular, the parameter fields marked by an asterisk are considered “required” by the user. All other parameters can be left empty, by assuming a default value (also reported in the hep page).



Fig. 48 – Example of a web page automatically open after the click on the help button

After completion of the parameter configuration, the “Submit” button launches the experiment.

After launch of an experiment, it can result in one of the following states:

- **Enqueued:** the execution is put in the job queue;
- **Running:** the experiment has been launched and it is running;
- **Failed:** the experiment has been stopped or concluded with any error occurred;
- **Ended:** the experiment has been successfully concluded;

Experiment Manager			
Experiment	Status	Last Access	Delete
myFirstExp	ended	2010-06-23	✖
myClassExp	running	2010-06-24	✖

Fig. 49 – Some different state of two concurrent experiments



DAta Mining & Exploration Program

3.7.1 Re-use of already trained networks

In the previous section a general description of experiment use cases has been reported. A specific more detailed information is required by the “Run” use case. As known this is the use case selected when a network (for example the MLP model) has been already trained (i.e. after training use case already executed).

The screenshot shows the 'Experiment Setup' window in DAMEWARE. At the top, there are two tabs: 'RESOURCE MANAGER' and 'Experiment Setup'. Below the tabs, there are two dropdown menus: 'Select a Functionality:' set to 'Classification_MLP' and 'Select a Running Mode:' set to 'Train'. A red circular 'HELP' button is visible. To the right, there are several input fields and dropdowns: 'Train Set*' (set to '/xor.csv'), 'Validation Set:', 'Network File:', 'number of input nodes*' (set to 2), 'number of nodes for hidden layer*' (set to 2), 'number of output nodes*' (set to 1), 'number of iterations:', 'error tolerance:', and 'training mode:'. A 'Submit' button is at the bottom right. A note '* = Field is Required' is displayed above the input fields.

Fig. 50 – An example of Classification_MLP training case for the XOR problem

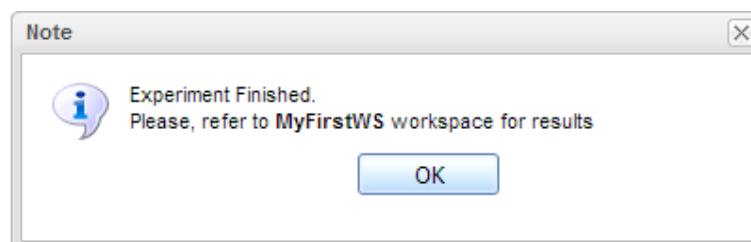


Fig. 51 – The popup status at the end of the XOR problem experiment

My Experiments				
Workspace: MyFirstWS				
Experiment	Status	Last Access	Delete	
xorTrain	ended	2013-09-09		
Download	AddInWS	File	Type	Description
		mlp_TRAIN_errorPlot.jpeg	jpeg	scatter plot of the epochs vs error
		mlp_TRAIN_error.csv	csv	epoch-error file
		mlp_TRAIN.log	txt	log file
		mlp_TRAIN_tmp_weights	mlp	net tmp file
		mlp_TRAIN_weights	mlp	trained network file

Fig. 52 – The list of output files after the XOR problem training experiment

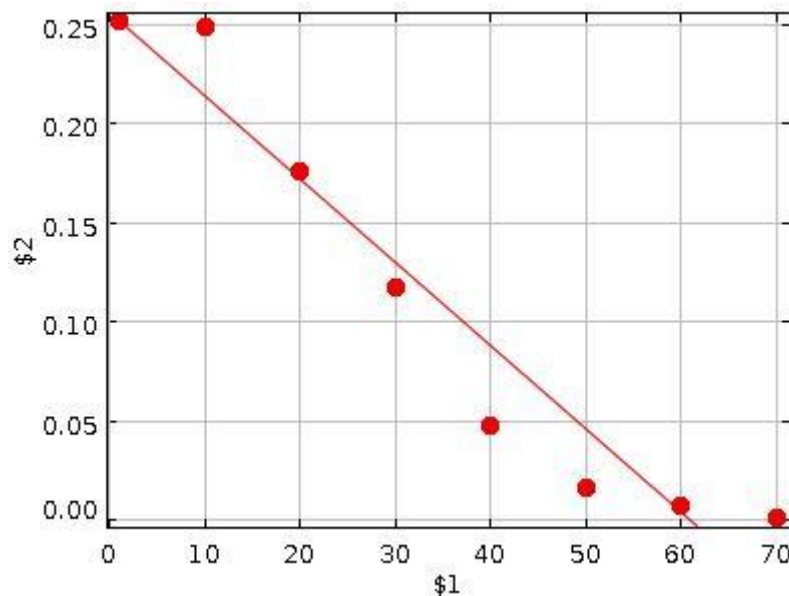


Fig. 53 – The training error scatter plot mlp_TRAIN_errorPlot.jpg downloaded from the experiment output list (x-axis is the training cycle, y-axis is the training mean square error)

The Run case is hence executed to perform scientific experiments on new data. Remember also that the input file does not include “target” values. The execution of a Run use case, for its nature, requires special steps in the DAME Suite. These are described in the following.

As first step, we require to have already performed a train case for any experiment, obtaining a list of output files (train or full use cases already executed). In particular in the output list of the train/full experiment there is the file *.mlp*. This file contains the final trained network, in terms of final updated weights of neuron layers, exactly as resulted at the end of the training phase. Depending on the training correctness this file has in practice to be submitted to the network as initial weight file, in order to perform test/run sessions on input data (without target values).

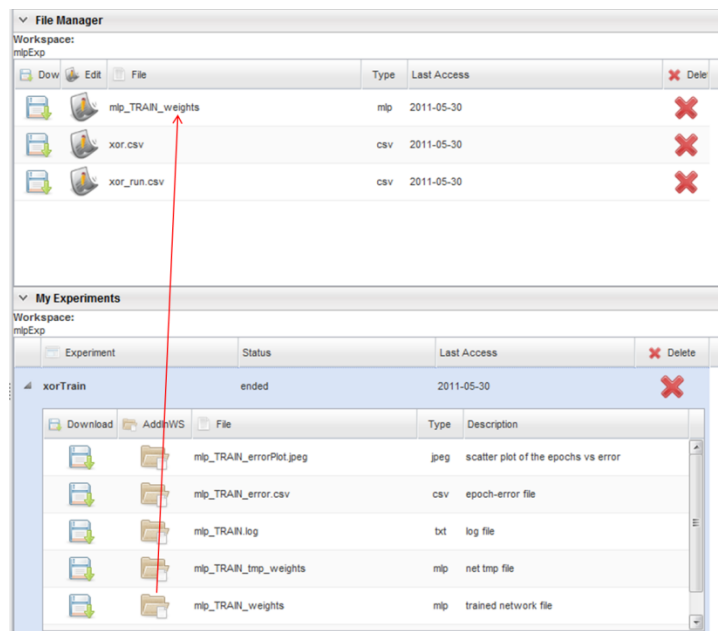


Fig. 54 – The operation to “move” the trained network file in the Workspace input file list

To do this, the output weight file must become an input file in the workspace file list, as already explained in section 3.5.4, otherwise it cannot be used as input of Test/Run use case experiment, Fig. 54. Also, the workspace currently active, hosting the experiment we are going to do, must contain a proper input file for Run cases, i.e. without target columns inside.

So far, the second step is to populate the workspace file list with trained network and Test/Run compliant input files and then to configure and execute the test experiment (see Fig. 55)

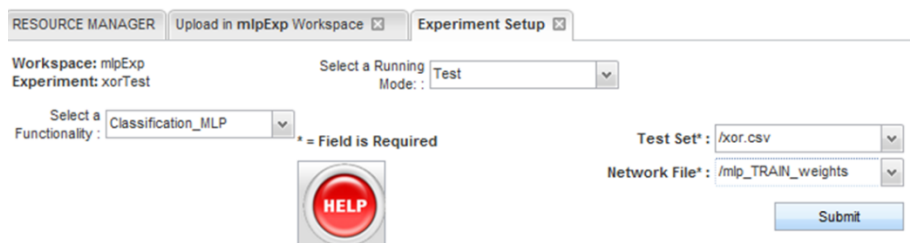


Fig. 55 –the configuration for the Run use case in the XOR problem



Data Mining & Exploration Program

Abbreviations & Acronyms

A & A	Meaning	A & A	Meaning
AI	Artificial Intelligence	IEEE	Institute of Electrical and Electronic Engineers
ANN	Artificial Neural Network	INAF	Istituto Nazionale di Astrofisica
ARFF	Attribute Relation File Format	JPEG	Joint Photographic Experts Group
ASCII	American Standard Code for Information Interchange	LAR	Layered Application Architecture
BoK	Base of Knowledge	MDS	Massive Data Sets
BP	Back Propagation	MLP	Multi Layer Perceptron
BLL	Business Logic Layer	MLPGA	MLP with Genetic Algorithms
CE	Cross Entropy	MLPQNA	MLP with Quasi Newton
CSOM	Clustering Self Organizing Maps	MSE	Mean Square Error
CSV	Comma Separated Values	NN	Neural Network
DAL	Data Access Layer	OAC	Osservatorio Astronomico di Capodimonte
DAME	Data Mining & Exploration	PC	Personal Computer
DAPL	Data Access & Process Layer	PI	Principal Investigator
DL	Data Layer	REDB	Registry & Database
DM	Data Mining	RIA	Rich Internet Application
DMM	Data Mining Model	SDSS	Sloan Digital Sky Survey
DMS	Data Mining Suite	SL	Service Layer
FITS	Flexible Image Transport System	SOFM	Self Organizing Feature Maps
FL	Frontend Layer	SOM	Self Organizing Maps
FW	FrameWork	SW	Software
GRID	Global Resource Information Database	UI	User Interface
GSOM	Gated Self Organizing Maps	URI	Uniform Resource Indicator
GUI	Graphical User Interface	VO	Virtual Observatory
HW	Hardware	XML	eXtensible Markup Language
KDD	Knowledge Discovery in Databases		

Tab. 2 – Abbreviations and acronyms



DAta Mining & Exploration Program

Reference & Applicable Documents

ID	Title / Code	Author	Date
R1	“The Use of Multiple Measurements in Taxonomic Problems”, in Annals of Eugenics, 7, p. 179--188	Ronald Fisher	1936
R2	<i>Neural Networks for Pattern Recognition</i> . Oxford University Press, GB	Bishop, C. M.	1995
R3	<i>Neural Computation</i>	Bishop, C. M., Svensen, M. & Williams, C. K. I.	1998
R4	Data Mining Introductory and Advanced Topics, Prentice-Hall	Dunham, M.	2002
R5	Mining the SDSS archive I. Photometric Redshifts in the Nearby Universe. <i>Astrophysical Journal</i> , Vol. 663, pp. 752-764	D’Abrusco, R. et al.	2007
R6	<i>The Fourth Paradigm</i> . Microsoft research, Redmond Washington, USA	Hey, T. et al.	2009
R7	Artificial Intelligence, A modern Approach. Second ed. (Prentice Hall)	Russell, S., Norvig, P.	2003
R8	Pattern Classification, A Wiley-Interscience Publication, New York: Wiley	Duda, R.O., Hart, P.E., Stork, D.G.	2001
R9	Neural Networks - A comprehensive Foundation, Second Edition, Prentice Hall	Haykin, S.,	1999
R10	<i>A practical application of simulated annealing to clustering</i> . Pattern Recognition 25(4): 401-412	Donald E. Brown D.E., Huntley, C. L.:	1991
R11	<i>Probabilistic connectionist approaches for the design of good communication codes</i> . Proc. of the IJCNN, Japan	Babu G. P., Murty M. N.	1993
R12	<i>Approximations by superpositions of sigmoidal functions</i> . Mathematics of Control, Signals, and Systems, 2:303–314, no. 4 pp. 303-314	Cybenko, G.	1989

Tab. 3 – Reference Documents



Data Mining & Exploration Program

ID	Title / Code	Author	Date
A1	SuiteDesign_VONEURAL-PDD-NA-0001-Rel2.0	DAME Working Group	15/10/2008
A2	project_plan_VONEURAL-PLA-NA-0001-Rel2.0	Brescia	19/02/2008
A3	statement_of_work_VONEURAL-SOW-NA-0001-Rel1.0	Brescia	30/05/2007
A4	MLP_user_manual_VONEURAL-MAN-NA-0001-Rel1.0	DAME Working Group	12/10/2007
A5	pipeline_test_VONEURAL-PRO-NA-0001-Rel.1.0	D'Abrusco	17/07/2007
A6	scientific_example_VONEURAL-PRO-NA-0002-Rel.1.1	D'Abrusco/Cavuoti	06/10/2007
A7	frontend_VONEURAL-SDD-NA-0004-Rel1.4	Manna	18/03/2009
A8	FW_VONEURAL-SDD-NA-0005-Rel2.0	Fiore	14/04/2010
A9	REDB_VONEURAL-SDD-NA-0006-Rel1.5	Nocella	29/03/2010
A10	driver_VONEURAL-SDD-NA-0007-Rel0.6	d'Angelo	03/06/2009
A11	dm-model_VONEURAL-SDD-NA-0008-Rel2.0	Cavuoti/Di Guido	22/03/2010
A12	ConfusionMatrixLib_VONEURAL-SPE-NA-0001-Rel1.0	Cavuoti	07/07/2007
A13	softmax_entropy_VONEURAL-SPE-NA-0004-Rel1.0	Skordovski	02/10/2007
A14	VONeuralMLP2.0_VONEURAL-SPE-NA-0007-Rel1.0	Skordovski	20/02/2008
A15	dm_model_VONEURAL-SRS-NA-0005-Rel0.4	Cavuoti	05/01/2009
A16	FANN_MLP_VONEURAL-TRE-NA-0011-Rel1.0	Skordovski, Laurino	30/11/2008
A17	DMPlugins_DAME-TRE-NA-0016-Rel0.3	Di Guido, Brescia, Cavuoti	14/04/2010
A18	BetaRelease_ReferenceGuide_DAME-MAN-NA-0009-Rel1.0	Brescia	28/10/2010
A19	BetaRelease_Model_MLP_UserManual_DAME-MAN-NA-0011-Rel1.0	Cavuoti, Brescia	30/11/2010
A20	BetaRelease_Model_SVM_UserManual_DAME-MAN-NA-0013-Rel1.0	Cavuoti, Brescia	30/11/2010
A21	BetaRelease_Model_MLPGA_UserManual_DAME-MAN-NA-0012-Rel1.0	Brescia	30/11/2010

Tab. 4 – Applicable Documents



DAta Mining & Exploration Program

Acknowledgments

The DAME program has been funded by the Italian Ministry of Foreign Affairs, the European project *VOTECH* (Virtual Observatory Technological Infrastructures) and by the Italian *PON-S.Co.P.E.* Leaders of the project are prof. G. Longo and prof. G.S. Djorgovski.

The current release of the data mining Suite is a miracle due mainly to the incredible effort of (in alphabetical order):

Giovanni Albano, Stefano Cavuoti, Giovanni d'Angelo, Alessandro Di Guido, Francesco Esposito, Pamela Esposito, Michelangelo Fiore, Mauro Garofalo, Marisa Guglielmo, Omar Laurino, Francesco Manna, Alfonso Nocella, Sandro Riccardi, Bojan Skordovski, Civita Vellucci

We want to really thank all actors who contribute and sustain our common efforts to make the whole DAME Program a reality, coming from University Federico II of Naples, INAF Astronomical Observatory of Capodimonte and Californian Institute of Technology.

Max

__oOo__



DAta Mining & Exploration Program



DAME Program
“we make science discovery happen”

