

## Use of Bioinformatics in Arrays

Peter Kalocsai and Soheil Shams

### 1. Introduction

Previous chapters have discussed in detail how to prepare, print and hybridize DNA arrays on various surfaces. Once hybridization is completed the next step is to scan in the glass, gel, or plastic slides with a specialized scanner to obtain digital images of the results of the experiment. The DNA expression levels are then quantified with the help of image-analysis software. After the image processing and analysis step is completed, we end up with a large number of quantified gene expression values. The data typically represent hundreds or thousands, in certain cases tens of thousands, of gene expression across multiple experiments. To make sense of this much information requires the use of various of visualization and statistical analysis techniques. One of the most typical microarray data analysis goal is to find statistically significant up- or downregulated genes; in other words outliers or "interestingly" behaving genes in the data. Other possible goals could be to find functional groupings of genes by discovering similarity or dissimilarity among gene-expression profiles, or predicting the biochemical and physiological pathways of previously uncharacterized genes.

Before any of that analysis could take place, one has to address the issues of normalization and possible transformation of the data so that, as much as possible, the quantified values would only represent true differences among gene expressions. On the other hand, finding a transformation that would make reading the data more understandable, or would change the distribution of the data in a way that it would be more suitable for certain statistical analysis could also be an important goal. In **Subheading 2.**, we discuss all these issues followed by a description of tools and algorithms implementing various

3. Before printing arrays: Slides work best when "aged" for 3 to 4 wk before use. Check that the polylysine coating is not opaque. Test print, hybridize, and scan sample slides to determine slide batch quality.
4. It is important to mark boundaries on the back of the slides because arrays become invisible after postprocessing.
5. Be sure the rack is bent slightly inward in the middle; otherwise, the slides may run into each other while shaking.
6. The amine groups on the Tris-HCl can react with the monofunctional NHS-ester of the Cy dye.
7. It is important to keep in mind that the DNA-binding curve for silica, on which the Qiagen PCR cleanup kit is based, is favorable at a low pH but falls off precipitously around pH 8.0. Thus, it is essential that the pH of your reaction be below 7.5 by the time it hits the Qiagen membrane.
8. If the array is exposed to air while the cover slip starts to fall off, you may see high background fluorescent signal on the side of the array.
9. Be careful not to allow Cy5 bleaching from scanners with strong lasers. If the arrays are not scanned immediately, they should be stored in light-tight slide boxes until they are scanned.

### References

1. Gillespie, D. and Spiegelman, S. A. (1965) A quantitative assay for DNA-RNA hybrids with DNA immobilized on a membrane. *J. Mol. Biol.* **12**, 829-842.
2. Ritossa, F., Malva, C., Boncinelli, E., Graziani, F., and Polito, L. (1971) The first steps of magnification of DNA complementary to ribosomal RNA in *Drosophila melanogaster*. *Proc. Nat. Acad. Sci. USA* **68**, 1580-1584.
3. Sinibaldi, R. M. and Cummings, M. R. (1981) Localization and characterization of rDNA in *Drosophila tumiditarsus*. *Chromosoma* **81**, 655-671.
4. Kafatos, F. C., Jones, C. W., and Estradiatis, A. (1979) Determination of nucleic acid sequence homologies and relative concentrations by a dot hybridization procedure. *Nucleic Acids Res.* **24**, 1541-1552.
5. Lockhart, D. J., Dong, H., Byrne, M. C., Follettie, M. T., Gallo, M. V., Chee, M. S., Mittmann, M., Wang, C., Kobayashi, M., Horton, H., and Brown, E. L. (1996) Expression monitoring by hybridization to high-density oligonucleotide arrays. *Nat. Biotechnol.* **14**, 1675-1680.
6. Schena, M., Shalon, D., Davis, R. W., and Brown, P. O. (1995) Quantitative monitoring of gene expression patterns with a complementary DNA microarray. *Science* **270**, 467-470.
7. Guo, Z., Guilfoyle, R. A., Thiel, A. J., Wang, R., and Smith, L. M. (1994) Direct fluorescence analysis of genetic polymorphisms by hybridization with oligonucleotide arrays on glass supports. *Nucleic Acids Res.* **22**, 5456-5465.

Several popular ways of normalizing and transforming microarray data are discussed as follows.

## 2.2. Normalization of the Data

One of the most popular ways to control for spotted DNA quantity and other slide characteristics is to do a type of local normalization by using two channels (red and green) in the experiment. For example, a cy5 (red)-labeled probe could be used as control prepared from a mixture of cDNA samples or from normal cDNA. Then a cy3 (green)-labeled experimental probe could be prepared from cDNA extracted from a tumor tissue. The normalized expression values for every gene would then be calculated as the ratio of experimental and control expression. This method can obviously eliminate a great portion of experimental variation by providing every spot (gene) in the experiment with its own control. Developing on these ideas, three-channel experiments are underway where one channel serves as control for the other two. In this case, the expression values of both experimental channels would be divided by the same control value.

In addition to the foregoing-described local normalization, global methods are also available in the form of "control" spots on the slide. Based on a set of these control spots, it becomes possible to control for global variation in overall slide quality or scanning differences.

These procedures describe some physical measures in terms of spotting characteristics that one can take to normalize the microarray data. However, even after the most careful two or more channel spotting with the use of control spots it is still possible that undesired experimental variation contaminates the expression data. On the other hand, it is also possible that all or some of these physical normalization techniques are missing from the experiment in which case it is even more important to find other ways of normalization. Fortunately, for both of these scenarios additional statistics based normalization methods are available to further clean up the data.

For example, the situation can happen when gene expression values in one experiment are consistently and significantly different from another experiment for the same set of genes due to quality differences between the slides or the printing or scanning process or possibly due to some other factor (see Fig. 1).

It might be very misleading to compare the expression values of the two files plotted in Fig. 1 without any normalization. When subtracting the mean value and dividing by the standard deviation in both experiments we can make a much more realistic comparison. The same data are plotted in Fig. 2 after normalization. Notice that most of the gene expressions now lay on the identity line as it would be expected for two experiments on the same set of genes.

We might mention that statistics based normalization or standardization can be accomplished in many different ways involving either the mean, median,

multivariate techniques that could help solve the data analysis and visualization needs of the DNA array research community in Subheading 3.

## 2. Data Normalization and Transformation

### 2.1. Examination of the Data

Before even starting to analyze the DNA array data, what one has to be aware of is that no matter how powerful statistical methods are used the analysis is still crucially dependent on the "cleanness" and distributional properties of the data. Essentially, the two questions to ask before starting any analysis are:

1. Does the variation in the data represent true variation in expression values or it is contaminated by differences in expression due to experimental variability?
2. Is the data "well-behaving" in terms of meeting the underlying assumptions of the statistical analysis techniques that are applied to it?

It is easy to appreciate the importance of the first point. The significance of the second one comes from the fact that most multivariate analysis techniques are based on underlying assumptions such as normality and homoscedasticity. If these assumptions are not met at least approximately, then the whole statistical analysis could be distorted and statistical tests might be invalid. Fortunately, there are a variety of statistical techniques available to help us answer "Yes" to the above two questions respectively:

1. Normalization (standardization).
2. Transformation.

Normalization, and as a special form of normalization standardization, can help us separate true variation from differences due to experimental variability. This step might be necessary, as it is quite possible that due to the complexity of creating, hybridizing, scanning, and quantifying microarrays variation originating from the experimental process contaminates the data. During a typical microarray experiment, many different variables and parameters can possibly change, hence differentially affect, the measured expression levels. Among these are slide quality, pin quality, amount of DNA spotted, accuracy of arraying device, dye characteristics, scanner quality, and quantification software characteristics, to name a few. The various methods of normalization aim at removing, or at least minimizing, expression differences due to variability in any of these types of conditions.

As will be discussed later, the various transformation methods all aim at changing the variance and distribution properties of the data in such a way so that it would be closer in meeting the underlying assumptions of the statistical techniques applied to it in the analysis phase. The most common requirements of statistical techniques are for the data to have homologous variance (homoscedasticity) and normal distribution (normality).



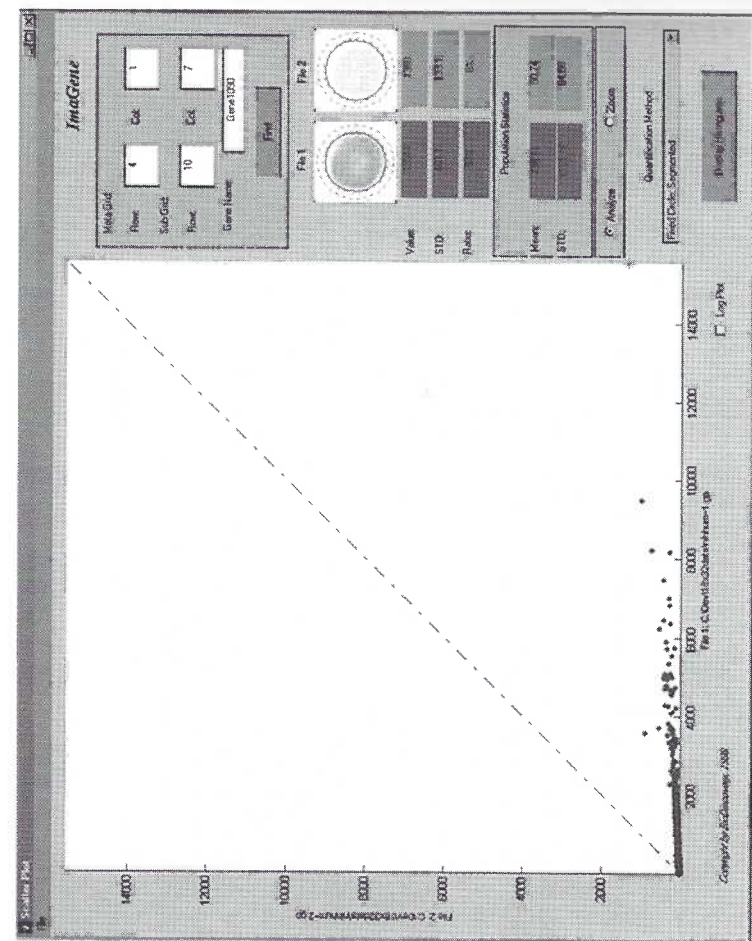


Fig. 1. Scatterplot of two experiments where the signal values are significantly stronger in one file than in the other. File 1 is the red value and File 2 the green. Reprinted with permission from **ref. 8**.

mode, or possibly some other statistics of the data. As a general rule, the choice of statistics should depend on the distributional properties of the data.

### 2.3. Transformation of the Data

Although there are many different data transforms available, the most frequently used procedure in the microarray literature is to take the logarithm of the quantified expression values (1,2). An often-cited reason for applying such a transform is to equalize variability in the typically wildly varying raw expression scores. If the expression value was calculated as a ratio of experimental over control conditions then an additional effect of the log-transform will be to equate up- and downregulation by the same amount in absolute value scores ( $\log_{10} 2 = 0.3$  and  $\log_{10} 0.5 = -0.3$ ). Another important side effect of the log transform is bringing the distribution of the data closer to normal. Having reasonable grounds of meeting the normality and homoscedasticity assump-

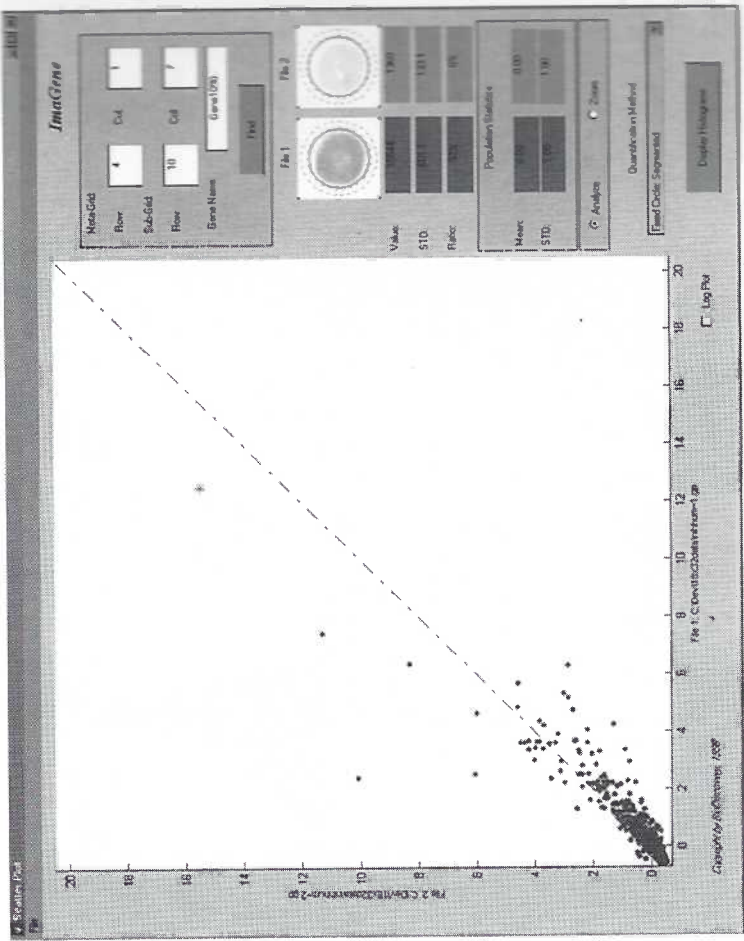


Fig. 2. The same data as in **Fig. 1** after normalization. File 1 is the red value and File 2 the green. Reprinted with permission from **ref. 8**.

tions after the log transform it is now much better justified to use a variety of parametric statistical analysis methods on the data.

As shown in **Fig. 3** (left panel), without the log transform we get a very peaked distribution of the data with a very long positive tail. This distribution is very far from normal, which violates the assumptions made by many standard parametric statistical analysis methods. The distribution of the log-transformed data in the right panel is visibly much closer to that of the normal distribution, which could also be verified by a simple normality test. Certainly, the number of transforms that could be experimented with is almost limitless, but the most frequently applied once are  $\log(x)$ ,  $\sqrt{x}$ ,  $1/x$ , and  $\arcsin \sqrt{x}$ . In fact, the choice of transformation should be dependent on which one brings the data closest to the requirements of homogeneity of variance and normal distribution (3). It turns out that for microarray expression data the usually applied log transform provides the best solution. On the other hand when nonparametric analysis techniques are used, which have no restrictions on the distributional

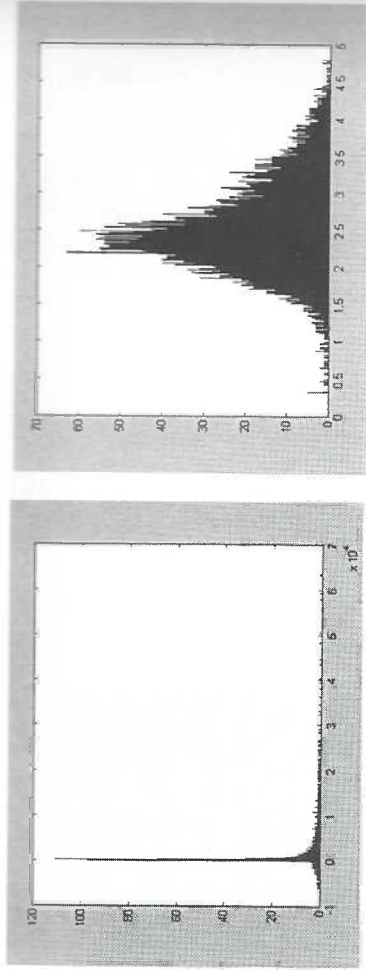


Fig. 3. Histograms of gene expression data involving 600 genes across 21 experiments. The x-axis indicates the expression values and the y-axis shows the number of genes with a particular gene expression level. The left and right panels show the data before and after the log transform, respectively. Reprinted with permission from ref. 8.

properties of the data, the transformation step might not be necessary at all. This flexibility of nonparametric methods comes at a price of usually requiring more data and, in many cases, providing less accuracy than parametric techniques.

### 3. Data Analysis

In Subheading 2, we gave an overview of the most important data preprocessing steps that one might consider before starting analyzing the data. In what follows we describe, with increasing complexity, the analysis techniques that appear to be the most useful for DNA array data.

#### 3.1. Scatterplot

Probably the simplest analysis tool for microarray data visualization is the scatterplot, as introduced earlier. In a scatterplot, each point represents the expression value of a gene in two experiments, one plotted on the x-axis and the other one on the y (see Fig. 2). In such a plot, genes with equal expression values would line up on the identity line (diagonal) with higher expression values further away from the origin. Points below the diagonal represent genes with higher expression in the experiment plotted on the x-axis. Similarly, points above the diagonal represent genes with higher expression values in the experiment plotted on the y-axis. The further the point away is from the identity line, the larger is the difference between its expression in one experiment compared with the other.

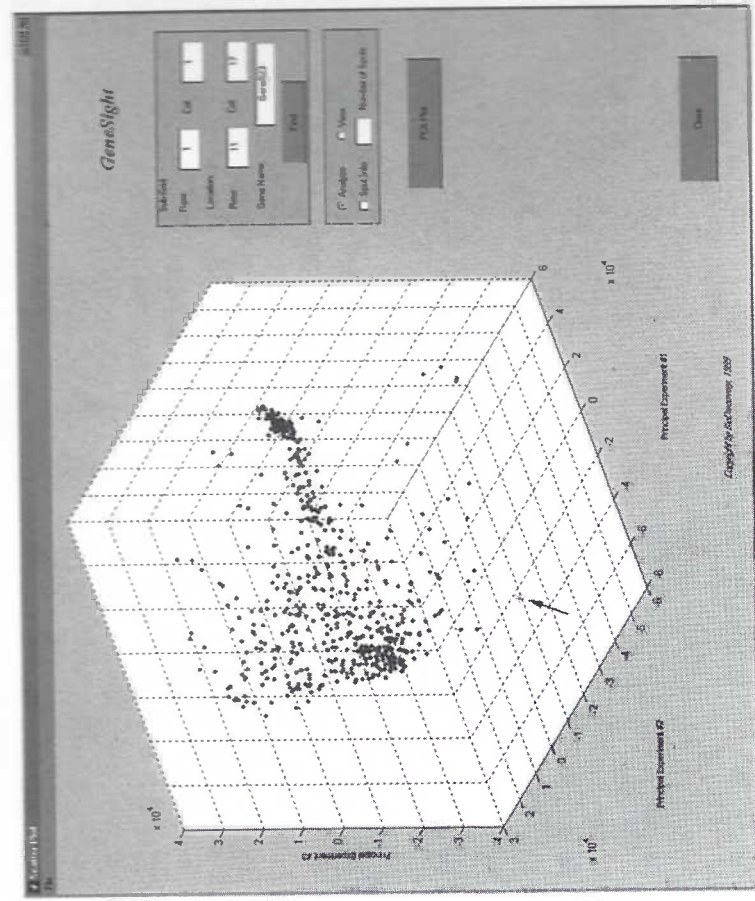


Fig. 4. Principal component analysis on 600 genes across 20 experiments. Gene scores are plotted on the first three principal components. The gene pointed to by the arrow shows a possible outlier.

### 3.2. Principal Component Analysis

It is easy to see how the scatterplot is an ideal tool for comparing the expression profile of genes in two experiments. Even three experiments could be plotted and compared in a 3-dimensional scatterplot. What can we do though when more than 3 experiments are to be analyzed and compared with each other? In case of 20 experiments, for example, we cannot draw a 20-dimensional plot. Fortunately, there are techniques available in statistics for dimensionality reduction, such as principal component analysis, which are able to compress the data into two or three dimensions (that we can plot) while preserving most or all the variance of the original dataset. Figure 4 is in fact a 3-dimensional plot of 600 genes in 21 experiments indicating the scores of all 600 genes on the first three principal components. Of course, lower-ranked principal components could also be plotted, three or less at a time, with the understanding that they account for less and less of the overall variance in the data.



This multivariate technique is frequently used to provide a compact representation of large amounts of data by finding the axes (principal components) on which the data varies the most. In principal component analysis, the coefficients for the variables are chosen such that the first component explains the maximal amount of variance in the data. The second principal component is perpendicular to the first one and explains maximum of the residual variance. The third component is perpendicular to the first two and explains maximum of the still remaining variance. This process is continued until all the variance in the data is explained. The linear combination of gene expression levels on the first three principal components could easily be visualized in a three-dimensional plot (*see Fig. 4*). This method, just like the scatterplot earlier, provides an easy way of finding outliers, in the data; genes that behave differently than most of the genes across a set of experiments. With a transpose of the data matrix, the experiments could also be plotted to find out possible groupings and/or outliers of experiments. Recent findings show that this method should be able to even detect moderate-sized alterations in gene expression (2). In general, principal component analysis provides a rather practical approach to data reduction, visualization, and identification of unusually behaving outlier genes and/or experiments.

### 3.3. Parallel Coordinate Planes

Two- and three-dimensional scatterplots and principal component analysis plots are ideal for detecting significantly up- or downregulated genes across a set of experiments. These methods do not provide an easy way of visualizing progression of gene expression over several experiments, however. These types of questions usually come up in time series experiments where gene expression is measured every two hours, for instance. The important question in this case is how gene expression progresses over the duration of the entire experiment. The parallel coordinate planes plotting technique is best suited to answer these types of questions. With this method experiments are ordered on the  $x$ -axis and expression values plotted on the  $y$ -axis. All genes in a given experiment are plotted at the same location on the  $x$ -axis; only their  $y$  location varies. Another experiment is plotted at another  $x$  location in the plane. Typically, the progression of time would be mapped into the  $x$ -axis by having higher  $x$  values for experiments done at a later time or vice versa. By connecting the expression values for the same genes in the different experiments one can obtain a very intuitive way of depicting the progression of gene expression (*see Fig. 5*).

Showing changes in expression pattern during the cell's life cycle can readily be visualized with parallel coordinate planes. Not only does this make this type of display very easy to follow the changes in expression level over time, but it could well be applied to any other type of data as well, such as comparison of

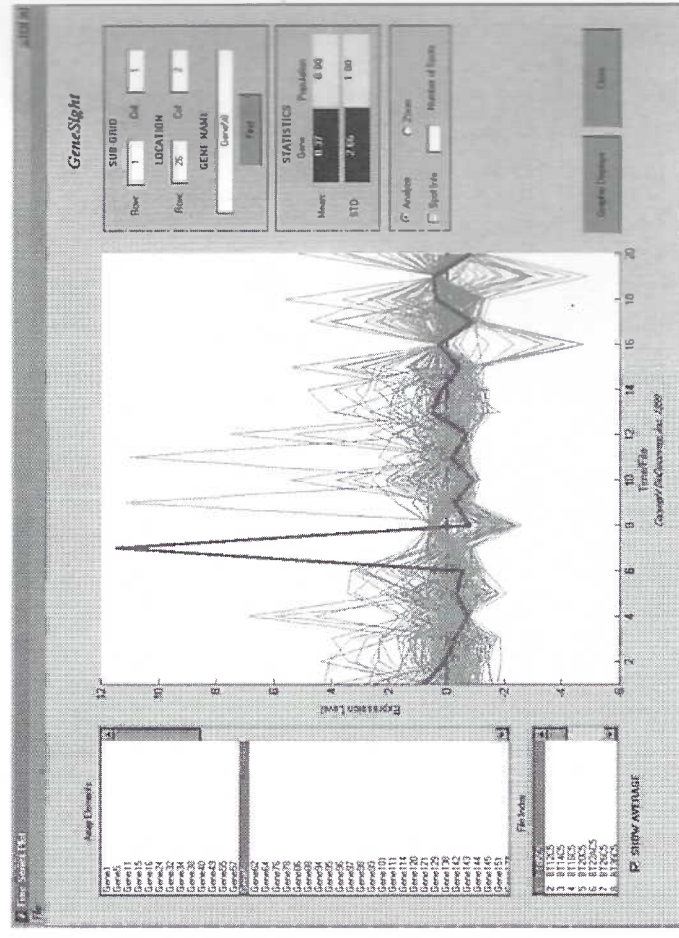


Fig. 5. The parallel coordinate plot displays the expression levels of all genes across all experiments/files in the analysis. On the  $x$ -axis the experiments or experimental files are plotted. The  $y$ -axis shows the expression level of all genes across all experiments.

expression level in different tissue types, for example. Due to the easy detection of unusual expression patterns, this type of plot can also be used well for outlier detection.

In addition, by applying different curve-fitting techniques any expression pattern over time could be searched for in the data. By adjusting the required closeness of fit the number of chosen expression patterns could also be controlled.

### 3.4. Cluster Analysis

Another frequently asked question related to microarrays is finding groups of genes with similar expression profiles across a number of experiments. The most often-used multivariate technique to find these groups is cluster analysis. Essentially, this method accomplishes the sorting of the data with grouping (clustering) genes with similar expression patterns closer to each other. This technique can help establish functional groupings of genes or predicting the biochemical and physiological pathways of previously uncharacterized genes.

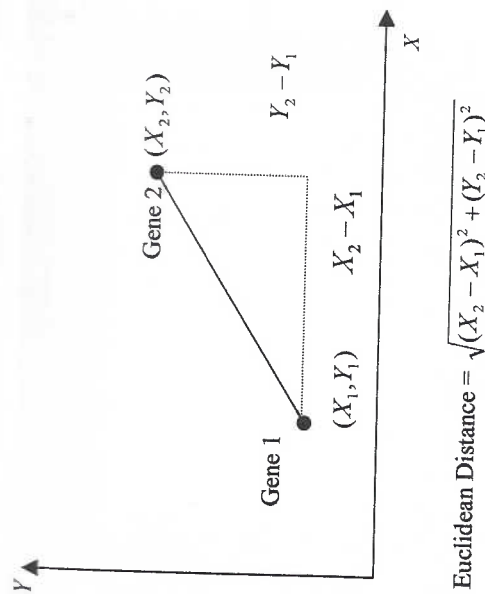


Fig. 6. Illustration of calculating the Euclidean distance between two genes in two experiments.

The clustering method that is most frequently used in the literature for finding groups in microarray data is hierarchical clustering (*1*). This method typically operates on a similarity or distance measure of the data such as correlation, Euclidean, squared Euclidean, or city-block (Manhattan) distance. To demonstrate the meaning of distance between two gene expressions values in two experiments, the calculation of the Euclidean distance is shown below in **Fig. 6**. The expression values for the two experiments are plotted on the  $X$ - and  $Y$ -axes, respectively. The expression values for a gene in the two experiments were and when plotted they produce the point "gene 1" in the  $X$ - $Y$  plane. The point "gene 2" is produced similarly. From the Pythagorean theorem, the Euclidean distance between points gene 1 and gene 2 can easily be calculated.

This example was two-dimensional, as there were only two experiments, but it can easily be extended to  $N$  dimensions, where  $N$  is the number of experiments in the study. Calculating the correlation matrix of the data is even less extensive computationally than the Euclidean distance, and in many situations it already produced quite sufficient results (*1,4*).

In addition to calculating the correlation or distance matrix, in most cases a linkage rule also has to be specified to indicate how distance should be calculated between groups and when groups are supposed to be joined together. The most popular linkage rules are single, complete, average linkage, or the centroid method. For example, in the average-linkage method, the distance between two clusters is calculated as the average distance between all pairs of objects in the two different clusters. As a result of the grouping process, a tree

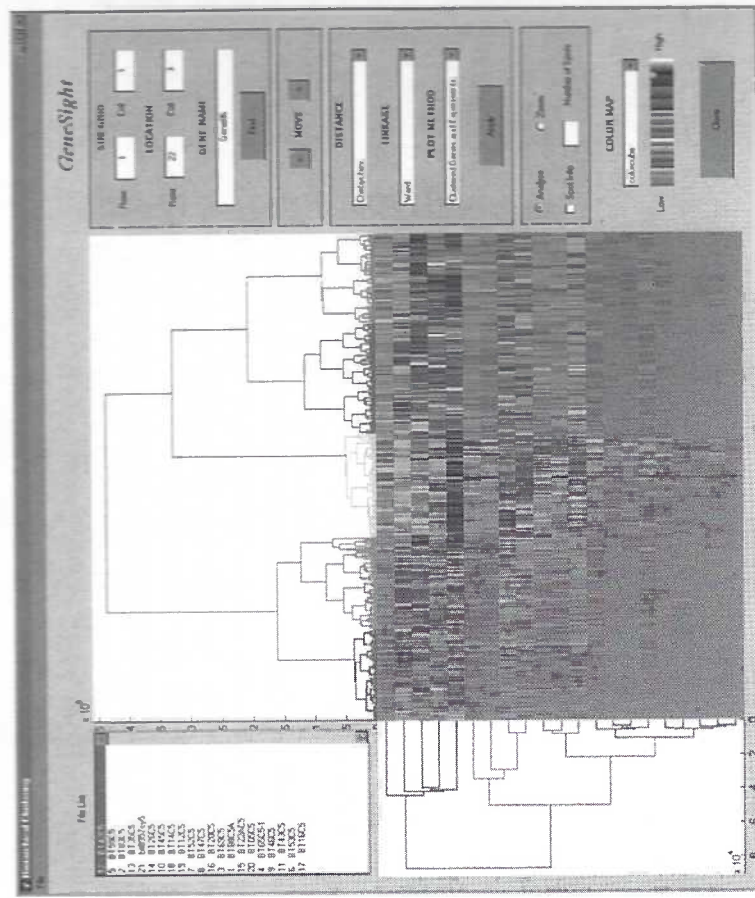


Fig. 7. Color-coded gene expression values for 600 genes (horizontally) in 21 experiments (vertically). Simultaneous clustering of genes and experiments is visualized by the top and left-side dendrograms, respectively.

of connectivity of observations emerges that can easily be visualized as dendrograms. For gene expression data, not only the grouping of genes but the grouping of experiments might also be important. When both are considered, it becomes easy to simultaneously search for patterns in gene expression profiles and across many different experimental conditions (see **Fig. 7**).

Every colored block in the middle panel of **Fig. 7** represents the expression value of a gene in an experiment. The 600 genes are plotted horizontally and the 21 experiments are plotted vertically. The color code is located in the lower-right corner. The dendrogram for genes is located just above the color-coded expression values with one arm connected to every gene in the study. As shown in **Fig. 7**, there are three major groups of genes in the study: The group on the left represents genes that have about medium-level expression in most experiments, but there also some experiments in the lower part with low expression. The middle, smaller, group represents genes that maintained medium-level



expression across all experiments with very high expressions in some of the experiments in the top part. The third group shows genes with low, medium, and high expression levels in certain experiments. These three major groups could then be subdivided further into smaller groups as shown in Fig. 7. The dendrogram for experiments is on the left showing the grouping of the 21 experiments in the study. The experiments with high overall expression levels are clustered together on the top. Experiments with medium-level expression are in the middle, whereas experiments with low-level expressions are grouped together in the bottom of the color-coded figure.

Although currently hierarchical clustering is an often employed way of finding groupings in the data, other nonhierarchical ( $k$  means) methods are likely to gain more popularity in the future with the rapidly growing amounts of data and the ever increasing average experiment size. A common characteristic of nonhierarchical approaches is to provide sufficient clustering without having to create the full distance or similarity matrix, but with minimizing the number of scans of the whole data set.

Cluster analysis is currently by far the most frequently used multivariate technique to analyze gene expression data. We have to emphasize, however, that it is also the simplest such method. Cluster analysis is typically employed when there is no *a priori* knowledge about the data available. We are at the very beginning of understanding the gene interaction network of even some of the simplest genomes, but it would certainly be misleading to say that nothing is known about the functionality of genes in certain genomes. For example, the Munich Information Center for Protein Sequences Yeast Genome Database classifies genes belonging into functional classes such as: tricarboxylic-acid pathway, respiration chain complexes, cytoplasmic ribosomes, and many others (5). Many of these functional categories represent genes, which, on biological grounds, are expected to have similar expression profiles across a set of experiments (1,4). One could, of course, apply the previously described clustering scheme to group genes with similar expression profiles and from the known genes in each group conclude which group represents which biological functionality. With such a procedure, one might find that the clustering actually came up with groups that biologically make sense, but the opposite is equally possible. Depending on the chosen algorithm, some of its parameters or just due to characteristics of the data it is also possible that the found clustering has no biological significance at all.

#### 4. Conclusions

In this chapter, we gave an overview of the most popular data analysis and visualization techniques that are used with microarray expression experiments. In our discussion, we started out with the simplest tools such as the scatterplot,

principal component analysis, and showing the data in parallel coordinate planes with gradually working our way toward the more complex analysis methods such as the various forms of clustering methods. The discussion should give the reader an overview of the currently used most popular analysis techniques with a peek about what to expect in the near future.

On a related note, no one would argue that even with the best possible algorithms the result of an analysis is crucially dependent on the quality of the data. Because of that importance, a whole section is committed to the various ways how one could improve data quality. That is, the numerous normalization and data transformation methods are discussed that have already proved useful for microarray data or at least have a chance of showing applicability in future analyses.

Currently, cluster analysis is the most popular multivariate technique that is used to find structure in microarray data. As pointed out earlier it is not without limitations, but as probably the simplest possible multivariate technique it quickly gained popularity. The authors predict that, as the field matures, we are likely to see a shift in the direction of more sophisticated classification methods appearing in the literature. Even though classification is certainly a more direct way of finding structure in the data than clustering, it still lacks the complexity that is required to capture all the connectivity and interdependence among genes in a genome.

Keep in mind that probably the ultimate goal of analyzing microarray data is, at some point, to discover how genes are related and affect each other, and dependent on one another. Accordingly, the modeling device describing this interdependency has to have a matching level of complexity. There are not too many modeling tools out there that would fit these requirements. Some of the possible candidates are multilayer neural networks, system of partial differential equations and structural equation modeling. At this point, it would be too early and also hard to tell which one or several of these and other methods will turn out to be the most applicable modeling tool(s), but with the rapidly accumulating expression data these techniques are bound to appear in the relatively near future. Some early examples of applying neural networks to explain gene data (6) and mapping out the connectivity pattern of smaller regulatory networks (7) are already available.

#### References

1. Eisen, M. B., Spellman, P. T., Brown, P. O., and Botstein D. (1998) Cluster analysis and display of genome-wide expression patterns. *Proc. Natl. Acad. Sci.* **95**, 14,863–14,868.
2. Hilsenbeck, S. G., Friedrichs, W. E., Schiff, R., O'Connell, P., Hansen, R. K., Osborne, C. K., and Fuqua, S. A. W. (1999) Statistical analysis of array expres-

- sion data as applied to the problem of tamoxifen resistance. *J. Natl. Cancer Inst.* **91**, 453-459.
3. Ferguson, G. A. and Takane Y. (1998) *Statistical Analysis in Psychology and Education*, McGraw-Hill, New York.
  4. Brown, P. O. and Botstein, D. (1999) Exploring the new world of the genome with DNA microarrays. *Nat. Genet. Suppl.* **21**, 33-37.
  5. Munich Yeast Genome Database (1999) Munich information center for protein sequences yeast genome database. <http://www.mips.biochem.mpg.de/proj/yeast>.
  6. Yuh, C., Bolouri, H., and Davidson, E. H. (1998) Genomic cis-regulatory logic: experimental and computational analysis of a sea urchin gene. *Science* **279**, 1896-1902.
  7. Tamayo, P., Slonim, D., Mesirov, J., Zhu Q., Kitareewan, S., Dmitrovsky, E., Lander, E. S., and Golub, T. R. (1999) Interpreting patterns of gene expression with self-organizing maps: Methods and application to hematopoietic differentiation. *Proc. Natl. Acad. Sci.* **96**, 2907-2912.
  8. Zhou, Y., Kalocsai, P., Chen, J., and Shams, S. (2000) Information processing issues and solutions associated with microarray technology, in *Microarray Biochip Technology*, BioTechniques Books Publication, (Skena, M., ed.), Eaton Publishing, Natick, MA.

## Confocal Scanning of Genetic Microarrays

Arthur E. Dixon and Savvas Damaskinos

### 1. Introduction

Confocal scanning laser microscopes are best known for making very high-resolution images of small three-dimensional (3D) specimens by recording a series of two-dimensional confocal slices of the specimen at different focus positions. Additionally, confocal microscopes are excellent for measuring weak fluorescence from areas adjacent to areas of very strong fluorescence, because the confocal pinhole rejects light from areas outside the focus spot. The resolution of a confocal microscope is also better than that of a nonconfocal instrument, an important property when submicron resolution is required. At first glance, none of these properties seems important for imaging genetic microarrays, whose features range from 200 to 10  $\mu$ m in size, and may cover the entire surface of a microscope slide in a two-dimensional array. However, we will find that the array of genetic material, when deposited on a weakly fluorescent substrate (like a glass slide) or in contact with an aqueous buffer solution acts like a 3D specimen as far as imaging is concerned. Confocal imaging helps to reject much of the background signal from the glass slide or the aqueous solution, and it is also useful for rejecting fluorescence from areas that surround the focal spot in the focal plane.

#### 1.1. Review of Confocal Microscopy

An excellent source for general information on confocal microscopy is the Handbook of Biological Confocal Microscopy (1). Figure 1 shows a basic confocal microscope. Light from a point source (often a pinhole illuminated by a focused laser beam) passes through a beam splitter, expands to fill a microscope objective, and is focused to a tiny volume (the focal point) inside the specimen (solid lines in Fig. 1). Light reflected (or emitted) from that point in

From: *Methods in Molecular Biology*, vol. 170: *DNA Arrays: Methods and Protocols*  
 Edited by: J. B. Rampal © Humana Press Inc., Totowa, NJ