

# Multivariate Statistical Analysis 1: Dimension Reduction Methods

**Presented by**

**Alex Shaw**

**Sydney Informatics Hub**

**Core Research Facilities**

**The University of Sydney**



# Acknowledging SIH



All University of Sydney resources are available to Sydney researchers **free of charge**. The use of the SIH services including the Artemis HPC and associated support and training warrants acknowledgement in any publications, conference proceedings or posters describing work facilitated by these services.

*The continued acknowledgment of the use of SIH facilities ensures the sustainability of our services.*

## **Suggested wording for use of workshops and workflows:**

*“The authors acknowledge the Statistical workshops and workflows provided by the Sydney Informatics Hub, a Core Research Facility of the University of Sydney.”*

# During the workshop



Ask **short questions** or clarifications during the workshop (either by Zoom chat or verbally). There will be breaks during the workshop for longer questions.



Slides with this **blackboard icon** are mainly for your reference, and the material will not be discussed during the workshop.



**Challenge questions** will be encountered throughout the workshop.



# Learning Objectives

## **This workshop:**

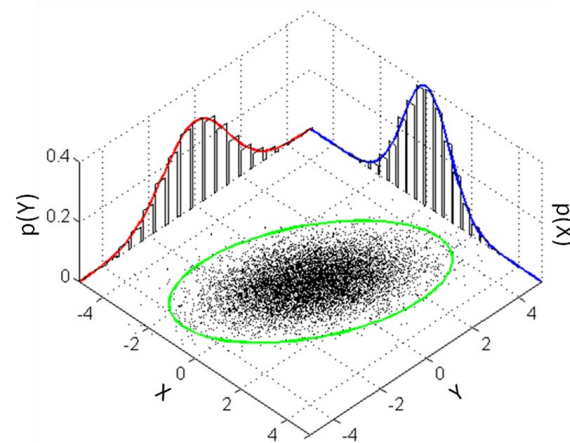
- Motivation for using multivariate methods
- Goals of multivariate methods
- A workflow for dimension reduction methods
- Introduction to dimension reduction methods:
  - Principal Components Analysis

## **Within accompanying software workflows:**

- Introduction to other dimension reduction methods:
  - Correspondence Analysis
  - nMDS

# Multivariate Statistics

- Multivariate statistics applied when you are considering the distribution of more than one variable (usually an outcome variable).
- We are usually interested in the relationships between such variables, i.e. how the variables vary together (their joint distribution). Otherwise, we could just perform [a simpler] analysis on each variable separately.
- It is common to have more than one variable in your study, but we often treat some variables as ‘fixed’ (i.e. not modelling their distribution). *E.g. in multiple linear regression, all the explanatory variables are fixed.*
- Multivariate statistics applies when you are modelling multiple variables in combination e.g. *in linear models with multiple outcome variables.*



[https://en.wikipedia.org/wiki/Multivariate\\_normal\\_distribution](https://en.wikipedia.org/wiki/Multivariate_normal_distribution)

# Why multivariate statistics?

- **Multivariate statistics** is a very broad topic with a huge array of different techniques. The motivation for using one of these techniques may include:
  - **Investigation of dependence:** relationships among variables
  - **Sorting and grouping:** identify relationships and groups among the entities/subjects under study
  - **Data reduction:** summarise multiple variables using a smaller set of ‘synthetic variables’
  - **Hypothesis testing:** group differences for a model with multiple outcome variables
  - **Prediction:** based on a multivariate model (multiple outcomes)

*Which of these has brought you to this workshop?*



# Special cases of multivariate statistics



- Multivariate statistics also come in to play in situations where we have repeated measurements (of the same thing) on the same subjects, especially longitudinal data
- Although these are indeed important applications of multivariate statistics, they are a special case and outside the scope of this workshop
- In this workshop, our examples will be limited to those where different things have been measured on the same individuals
- See our Linear Models series of workshops and/or book a consult with us to discuss analysis of repeated measures data

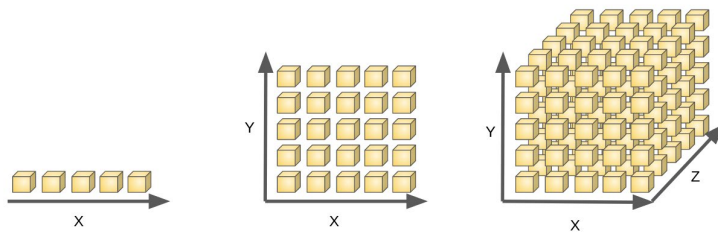
# Dimensionality Reduction Techniques



# Why do we need dimension reduction?

- To deal with multivariate data, in particular **high dimensional** multivariate data:
- Situations where we have as many (or more) observations/measurements of each subject as we have subjects. **Sometimes referred to as the ‘curse of dimensionality’**
- Common in fields that collect a lot of measurements: high throughput biological assays e.g. gene expression, neuroimaging

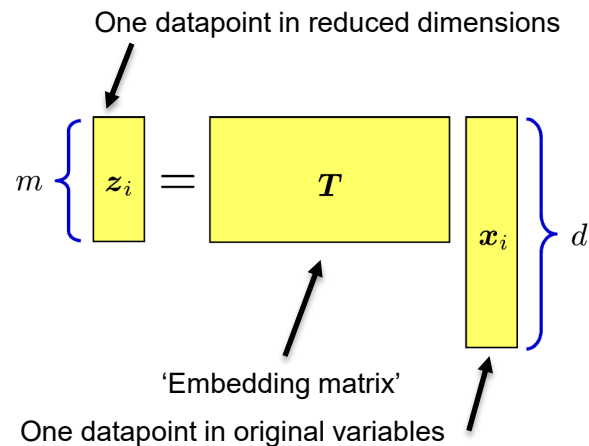
*“Even if the number of collected data points is large, they remain sparsely submerged in a voluminous high-dimensional space that is practically impossible to explore exhaustively”*



Nguyen LH, Holmes S (2019) Ten quick tips for effective dimensionality reduction. PLOS Computational Biology 15(6): e1006907. <https://doi.org/10.1371/journal.pcbi.1006907>

# Dimension Reduction (DR): what is it?

- You can think of dimension reduction as using the relationships between variables or subjects to concentrate information dispersed throughout your original dataset into a lower number of dimensions.
- This makes these relationships more interpretable and/or more analysis-ready
- Every datapoint in the space of original variables can be transformed to a point on in the reduced dimensional space, which facilitates exploration of your 'sparse' data



# Why do we need dimension reduction?

## To explain patterns in our data:

- We want to understand the relationships between variables (dependence) in situations where we have a lot of variables:
  - Are there patterns in associations/correlations involving more than two variables?
  - Do these patterns yield insight into **latent variables**, unobserved variables that represent some meaningful concept that can't be observed directly (e.g. personality traits, stage of disease progression)
- We want to understand the relationships between subjects:
  - Are there groups (**clusters?**) of subjects
  - How similar are our subjects to each other? Do they fall along a spectrum that correlates with some other variable that can be directly measured (e.g. different field sites w/ an environmental variable like soil PH, temperature, sunlight)

# Why do we need dimension reduction?

## To be able to analyse our data:

- When we have a lot of different measurements, we can't use standard techniques like linear models:
  - We have too many predictors to be able to fit a model and we may also have issues with:
  - Multicollinearity – situations where predictors are very similar to each other
    - Consequence: Difficulty accurately estimating effect of each predictor on the outcome. See **Model Building Workshop**
- We also can't complete the usual exploratory data analysis (EDA)
  - Too many measurements to be able to: look at all possible pairs of measurements or subjects to explore their relationship, detect outlier samples
- Sometimes, we just want to have fewer variables going into our [downstream] analysis:
  - Is there a way to summarise all the measurements into a few measurements (e.g. what are groups of genes that exhibit a similar response to drug treatment?)

# Dimension Reduction in the General Research Workflow

1. **Hypothesis Generation (Research/Desktop Review)**
2. Experimental and Analytical Design (sampling, power, ethics approval)
3. Collect/Store Data
4. Data cleaning
5. Exploratory Data Analysis (EDA)
6. Data Analysis aka inferential analysis
7. Predictive modelling
8. Publication

Discovery projects to generate new hypotheses e.g. new cell populations for further characterisation from single-cell expression data



# Dimension Reduction in the General Research Workflow

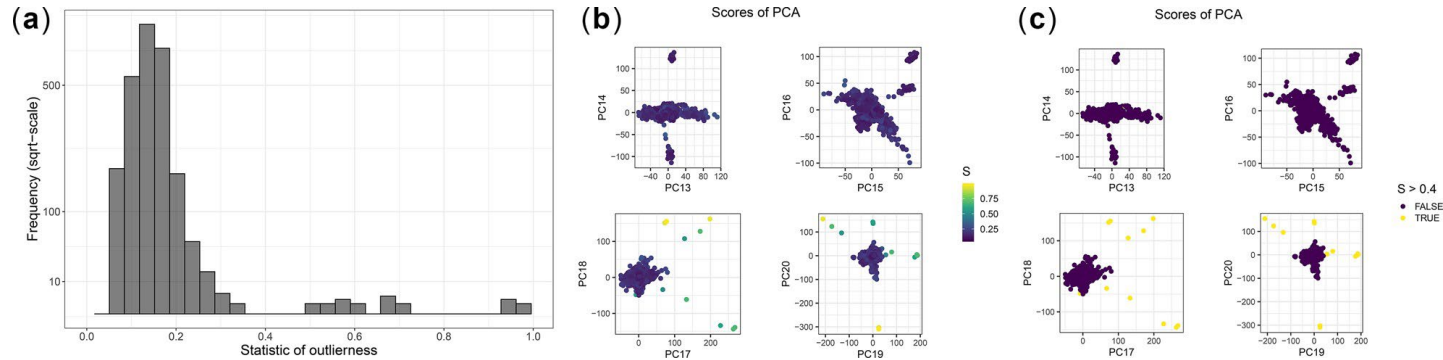
1. Hypothesis Generation (Research/Desktop Review)
2. Experimental and Analytical Design (sampling, power, ethics approval)
3. Collect/Store Data
4. **Data cleaning**
5. Exploratory Data Analysis (EDA)
6. Data Analysis aka inferential analysis
7. Predictive modelling
8. Publication

Multivariate techniques are an important part of QC (Quality Control) in high-throughput biology





**Fig. 2.** Outlier detection in the 1000 Genomes (1000G) project, using prob\_dist (Section 3.4)



*Bioinformatics*, Volume 36, Issue 16, 15 August 2020, Pages 4449–4457, <https://doi.org/10.1093/bioinformatics/btaa520>

The content of this slide may be subject to copyright: please see the slide notes for details.

# Dimension Reduction in the General Research Workflow

1. Hypothesis Generation (Research/Desktop Review)
2. Experimental and Analytical Design (sampling, power, ethics approval)
3. Collect/Store Data
4. Data cleaning
5. **Exploratory Data Analysis (EDA)**
6. Data Analysis aka inferential analysis
7. Predictive modelling
8. Publication

An important component of the research workflow. When there are multiple predictors, characterising the higher-order relationships between them is useful. See **Model Building Workshop**.

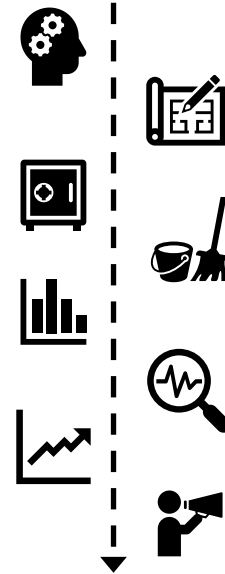




# Dimension Reduction in the General Research Workflow

1. Hypothesis Generation (Research/Desktop Review)
2. Experimental and Analytical Design (sampling, power, ethics approval)
3. Collect/Store Data
4. Data cleaning
5. Exploratory Data Analysis (EDA)
6. **Data Analysis aka inferential analysis**
7. **Predictive modelling**
8. Publication

PCA regression: perform linear regression in the presence of a high level of multicollinearity (see **Model Building Workshop**)



# What do we get out of dimension reduction?

## What are the outputs of this process:

- Synthetic Variables: We reduce the number of variables in our analysis by replacing original variables with a smaller number of synthetic variables. By 'synthetic variable', I mean one that *synthesises* information from multiple 'original variables'
- **Body Mass Index (BMI)** is a simple example of a synthetic variable:

$$\frac{\text{Weight in kg}}{(\text{Height in m})^2}$$

- Dimension reduction techniques combine information across variables\* to create useful synthetic variables, these may or may not be interpretable as **latent variables**

\*For simplicity and consistency I will mostly call the original variables just 'variables' and the synthetic variables 'dimensions' (or components)

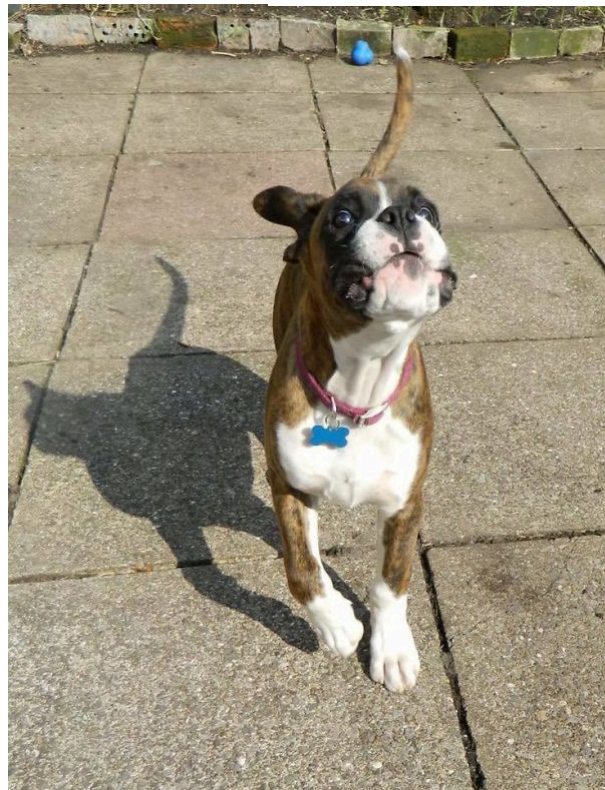
# What do we get out of dimension reduction?

- How do we get these outputs, and make sense of them? Let's start at the beginning:
- How do we explore a higher-dimensional dataset?
- Make a lower-dimensional projection plot
- Projection is an important geometrical concept to understand in dimension reduction. It is how we go from our high dimensional data to a more interpretable and explorable lower dimensional representation.

# Projections in lower dimensions

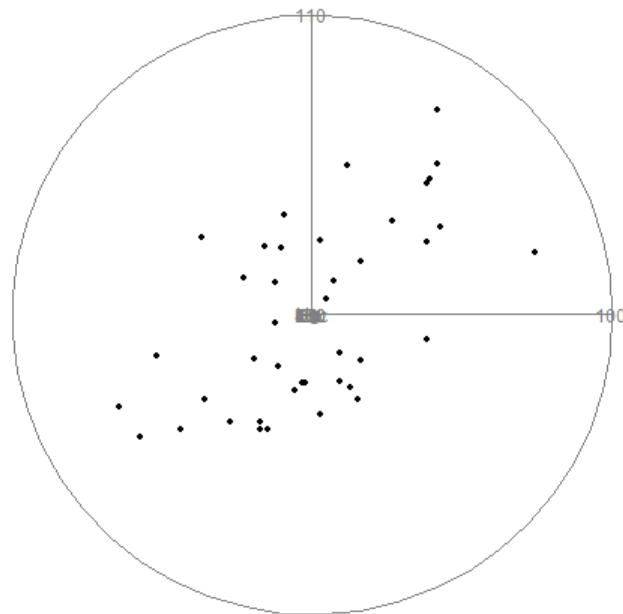


- We can't see in more than 3 dimensions, but if you have a 3D space we can visualise what projections of different vectors look like onto a 2D plane.
- The length of some vectors (D, E) is close to the length of their projections, while for others (C), the projection length is much shorter. This depends on the size of the vector and the angle between the vector and the plane
- The closer the actual vector length is to the length of its projection, the more accurately it is being represented by its projection



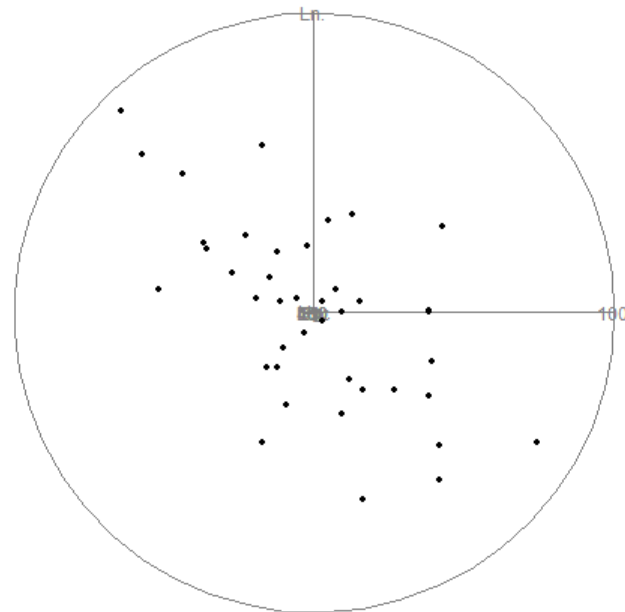
# Taking a tour(r)

- We can take a photo/project an image from any angle in higher dimensional space
- Some angles will be more informative than others regarding:
  - Original variable correlations
  - Clustering, outliers
- Let's start with choosing angles that compare 2 variables at a time in a set of decathlon performances (we will use PCA on this example dataset soon)



# Taking a grand tour

- Now let's just choose a path through the higher dimensional space at random
- We start to see relationships between variables and where each subject is, but how do we choose the most informative angle?



# A workflow for dimension reduction

# What is a workflow?

- Every statistical analysis is different, but all follow similar paths. It can be useful to know what these paths are
- We have developed practical, step-by-step instructions that we call ‘workflows’, that you can follow and apply to your research
- We have a general research workflow that you can follow from hypothesis generation to publication
- And statistical workflows that focus on each major step along the way (e.g. experimental design, power calculation, model building, analysis using linear models/survival/multivariate/survey methods)





# Statistical Workflows

- Our **statistical workflows** can be found within our workshop slides
- **Statistical workflows** are software agnostic, in that they can be applied using any statistical software
- **Software workflows** that show you how to perform the statistical workflow using particular software packages (e.g. R or Python) are forthcoming.
  - R software workflow is complete (but still evolving) for this workshop
  - There are other SIH hands-on workshops that show you how to perform dimension reduction in R or Python.



# A workflow for dimension reduction

0. Identify your variable types and perform appropriate Exploratory Data Analysis
1. Choose a dimension reduction method and run dimension reduction analysis
2. Examine the relationships between variables
3. Examine the relationships between subjects
4. Further summarising/interpretation. Choose how many dimensions to keep/examine.
5. [Optional] Downstream analysis

# Step 0 – Exploratory Data Analysis

- EDA is an important step in our general research workflow
  - See the EDA step in our Research Essentials workshop
  - Useful for choosing the right Dimension Reduction method
- Once you have chosen, there is some further EDA required – partly from the output of the method:
  - Want to make sure that the method is appropriate
  - Want to make sure that we use the right settings for our data and analysis questions

# Step 1 – Run Dimension Reduction Analysis

- The purpose of this workshop is to introduce the theory behind dimension reduction and show some basic examples
- You may have a single dimension reduction method you want to use, or you may want to run multiple methods and compare outputs
- Dimension reduction analysis can be performed in several statistical software packages
- If you are planning to use R, consult the software workflow for this workshop first
- SIH now offers two machine learning workshops series that include learning to run PCA either:
  - In R (using the Tidymodels approach)
  - In Python
  - See the [SIH Training Calendar](#) for how to dates and to enrol in these workshops



# Dimension reduction methods

| Method           | Input Data             | Method Class | Nonlinear | Complexity                      |
|------------------|------------------------|--------------|-----------|---------------------------------|
| → PCA            | continuous data        | unsupervised |           | $\mathcal{O}(\max(n^2p, np^2))$ |
| → CA             | categorical data       | unsupervised |           | $\mathcal{O}(\max(n^2p, np^2))$ |
| MCA              | categorical data       | unsupervised |           | $\mathcal{O}(\max(n^2p, np^2))$ |
| PCoA (cMDS)      | distance matrix        | unsupervised |           | $\mathcal{O}(n^2p)$             |
| → NMDS           | distance matrix        | unsupervised |           | $\mathcal{O}(n^2h)$             |
| Isomap           | continuous*            | unsupervised | ✓         | $\mathcal{O}(n^2(p + \log n))$  |
| Diffusion Map    | continuous*            | unsupervised | ✓         | $\mathcal{O}(n^2p)$             |
| Kernel PCA       | continuous*            | unsupervised | ✓         | $\mathcal{O}(n^2p)$             |
| t-SNE            | continuous/distance    | unsupervised | ✓         | $\mathcal{O}(n^2p + n^2h)$      |
| Barnes-Hut t-SNE | continuous/distance    | unsupervised | ✓         | $\mathcal{O}(nh \log n)$        |
| LDA              | continuous (X and Y)   | supervised   |           | $\mathcal{O}(np^2 + p^3)$       |
| PLS (NIPALS)     | continuous (X and Y)   | supervised   |           | $\mathcal{O}(npd)$              |
| NCA              | distance matrix        | supervised   | ✓         | $\mathcal{O}(n^2h)$             |
| Bottleneck NN    | continuous/categorical | supervised   | ✓         | $\mathcal{O}(nph)$              |
| STATIS           | continuous             | multidomain  |           | $\mathcal{O}(n^2P, nP^2)$       |
| DiSTATIS         | distance matrix        | multidomain  |           | $\mathcal{O}(n^2P, nP^2)$       |

Mentioned in this workshop



Tutorial example in software workflow

Nguyen LH, Holmes S (2019) Ten quick tips for effective dimensionality reduction. PLOS Computational Biology 15(6): e1006907.

<https://doi.org/10.1371/journal.pcbi.1006907>

<https://journals.plos.org/ploscompbiol/article?id=10.1371/journal.pcbi.1006907>

**Table 2. Example implementations.**

| Method             | R function                         | Python function                                                       |
|--------------------|------------------------------------|-----------------------------------------------------------------------|
| PCA                | <code>stats::prcomp</code>         | <code>sklearn.decomposition.PCA</code>                                |
| CATPCA             | <code>gifi::princals</code>        |                                                                       |
| CA                 | <code>FactoMineR::CA</code>        |                                                                       |
| MCA                | <code>FactoMineR::MCA</code>       |                                                                       |
| PCoA (cMDS)        | <code>stats::cmdscale</code>       | <code>sklearn.manifold.MDS</code>                                     |
| NMDS               | <code>ecodist::nmms</code>         | <code>sklearn.manifold.MDS</code>                                     |
| Isomap             | <code>vegan::isomap</code>         | <code>sklearn.manifold.Isomap</code>                                  |
| Diffusion Map      | <code>diffusionMap::diffuse</code> |                                                                       |
| (Barnes–Hut) t-SNE | <code>Rtsne::Rtsne</code>          | <code>sklearn.manifold.TSNE</code>                                    |
| LDA                | <code>MASS::lda</code>             | <code>sklearn.discriminant_analysis.LinearDiscriminantAnalysis</code> |
| PLS (NIPALS)       | <code>mixOmics::pls</code>         | <code>sklearn.cross_decomposition.PLSRegression</code>                |
| DiSTATIS           | <code>DistatisR::distatis</code>   |                                                                       |
| Procrustes         | <code>vegan::procrustes</code>     | <code>scipy.spatial.procrustes</code>                                 |

Software packages and function performing specified DR techniques available in R and python. R implementations are given as `package_name::function_name`; listed python functions come from `sklearn` and `scipy` libraries. The outputs of most linear DR methods can be visualized in R with `factoextra` package [25], used to generate a number of the plots in this article. Abbreviations: CA, correspondence analysis; CATPCA, categorical PCA; cMDS, classical multidimensional scaling; DR, dimensionality reduction; LDA, linear discriminant analysis; MCA, multiple CA; NIPALS, nonlinear iterative partial least squares; NMDS, nonmetric multidimensional scaling; PCA, principal component analysis; PCoA, principal CA; t-SNE, t-Stochastic Neighbor Embedding; PLS, partial least squares

<https://doi.org/10.1371/journal.pcbi.1006907.t002>

Nguyen LH, Holmes S (2019) Ten quick tips for effective dimensionality reduction. PLOS Computational Biology 15(6): e1006907.

<https://doi.org/10.1371/journal.pcbi.1006907>

<https://journals.plos.org/ploscompbiol/article?id=10.1371/journal.pcbi.1006907>

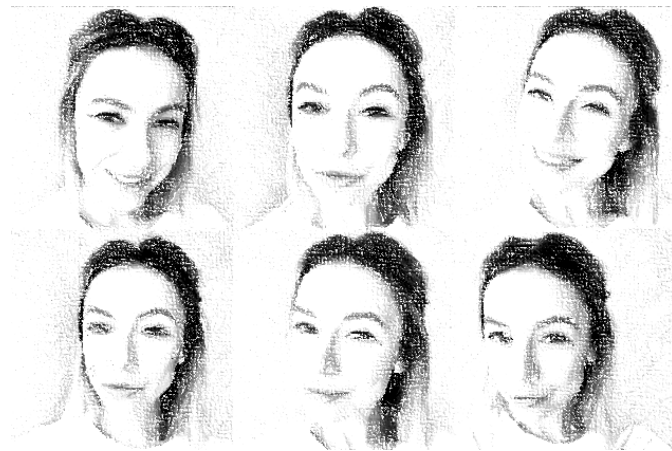
# Dimensionality Reduction Techniques: PCA-like

# Principal Component Analysis (PCA)



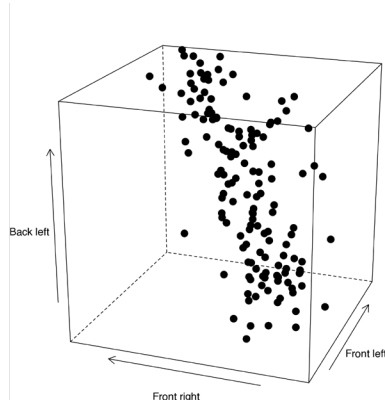
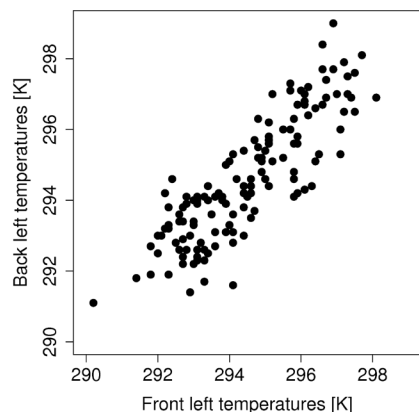
## Principal Components Analysis (PCA): Ever taken a selfie?

- There are many sites that advise on how to take the best selfie
- One element that is clearly important is, which angle? You're reducing from a 3D person to a 2D image so you are losing information
- We might choose the angle to highlight certain features of our face
- PCA is like taking a photo. PCA chooses an angle in higher-dimensional space that captures as much information as possible in a low number of dimensions.



# Mechanics of Dimension Reduction

- Recall our recent ‘tour’ through higher dimensional space. Think of your high-dimensional data as a ‘cloud’ of points in space. Each measured variable defines an axis, and each observation is located at a point in space defined by its measurements



*Returning to our question on which angle is most informative:*

*Can we describe these points in fewer dimensions while still retaining as much as possible of the shape of this point cloud?*

<https://learnche.org/pid/latent-variable-modelling/principal-component-analysis/visualizing-multivariate-data>



# Challenge question: Dimension reduction

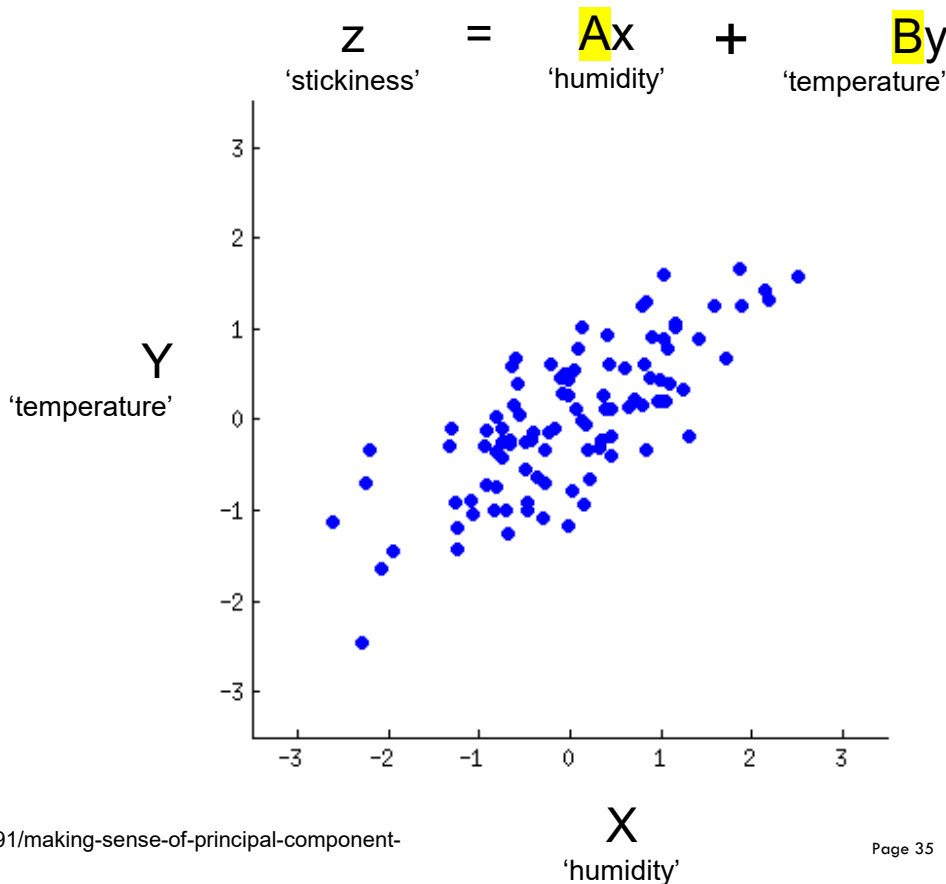
The scatter plot on the right represents measurements on two variables ( $x$  and  $y$ ). For illustrative purposes, let's say that  $x$  is temperature and  $y$  is humidity at midday for some set of days. The variables have been mean centred, so each measurement is +ve if above the mean and -ve if below.

We want to create a new scale  $z$  that measures 'stickiness', **some combination** of  $x$  and  $y$  that will preserve the most information, so that points that had relatively high measurements on both original variables have higher numbers in the new scale.

Draw a line through the points that will form a new axis (or measurement scale). All points will 'fall on to the line' by following a perpendicular path to this new axis.

Is this line of best fit, like that in a simple linear regression?

No. See: <https://shankarmysy.github.io/posts/pca-vs-lr.html>



# The PCA approach

- PCA chooses an angle in higher dimensional space, with a series of ‘principal directions’. These directions are part of the chosen angle/hyperplane (higher-dimensional plane)
- Projection of your original data points along these directions form Principal Components (PCs) for which PCA is named
- How are the principal directions/hyperplane chosen?
  - Maximises variation in the initial principal component (PC), PC1
  - All subsequent principal components are orthogonal/uncorrelated, while maximising variation sequentially (i.e.  $PC1_{\text{variation}} > PC2_{\text{variation}} > PC3_{\text{variation}} \dots$ )

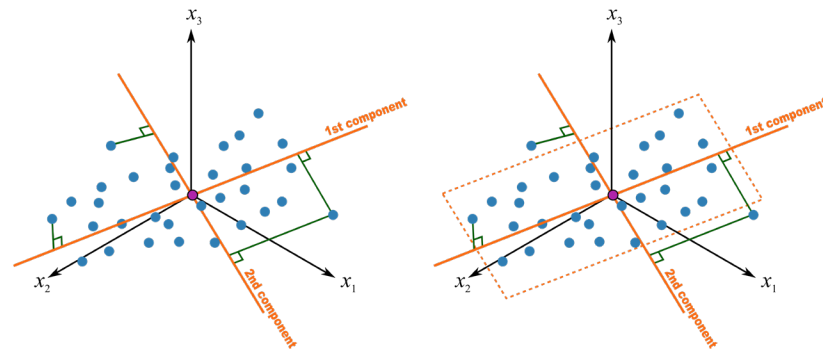
# What are the outputs from a PCA?

- Mathematically, each principal component is a linear combination or weighted sum of the original variables. The coefficients/weights are called the **loadings**
  - E.g.  $PC1 = Ax + By + Cz \dots$  where A, B and C are scalar coefficients/weights and x, y and z are points in the original coordinate system, i.e. the measurements for each subject
- The **loadings** (coefficients)\* allow us to calculate the projection of our data from the original variables to a **PC score** for each 'subject' on each individual principal component
- The PCA outputs are mathematically related to your covariance/correlation matrix
  - The principal directions are **eigenvectors** of the covariance/correlation matrix
  - The corresponding **eigenvalues** quantify the variation captured in the relevant principal component
  - The loadings are the **eigenvectors** that have been scaled by the [square root of the] **eigenvalues**

• Technical note: Different R packages provide either unscaled (eigenvectors), or scaled (loadings) directions as outputs. Terminology can vary, see <https://stats.stackexchange.com/questions/143905/loadings-vs-eigenvectors-in-pca-when-to-use-one-or-another>

# Projections in lower dimensions

- In a (PCA-style) loading plot:
  - The vectors represent original variables
  - The orientation of the plane is part of the PCA solution (two principal directions)
  - Projection of the variables (shown as vectors) shows us the **loadings plot**
  - Subjects can also be projected onto the plane (shown as points) **subjects plot**
  - Or both variables and subjects **biplot**



# An example dataset: Decathlon

- Some of the concepts are best explained using a concrete example
- Let's use the **decathlon** data set, which contains the performance of athletes in the 10 events of the decathlon
  - Track events: seconds to complete distance
  - Field events: distance travelled in metres
- So we have data from 41 performances across 10 variables
- Unless stated, all subsequent outputs are from this example

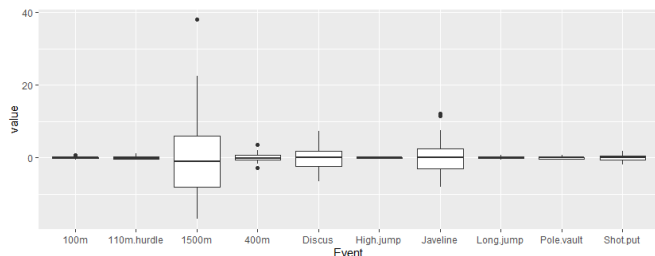
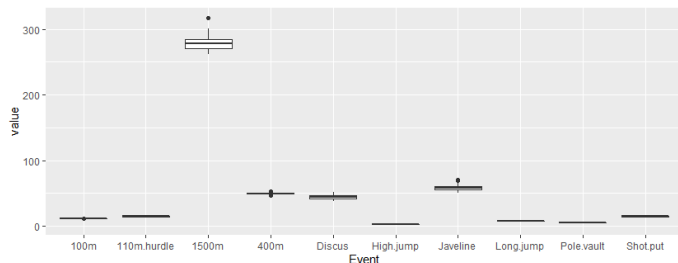


# A workflow for dimension reduction

0. **Identify your variable types and perform appropriate Exploratory Data Analysis**
  1. Run dimension reduction analysis
  2. Examine the relationships between variables
  3. Examine the relationships between subjects
  4. Further summarising/interpretation. Choose how many dimensions to keep/examine.
  5. Downstream analysis

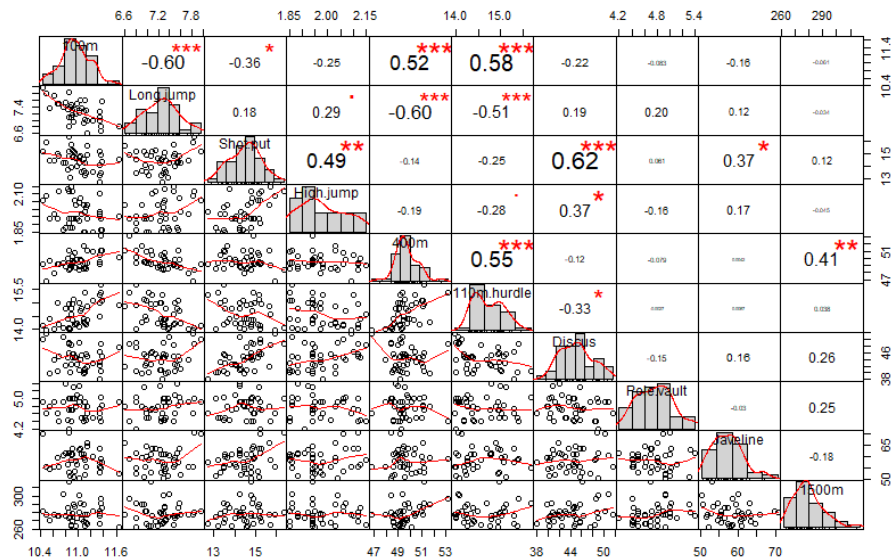


# Decathlon Dataset. Step 0: EDA



- We always mean centre the data (subtract the mean value from each variable) before PCA. Remember that we are using correlation, so the mean must be subtracted otherwise biases are introduced.
- It is also clear in this example that the different variables are on a different scale of measurement (different length track events with time in s, field events are distances in m)
- Beyond this, the variances for our variables are quite different to each other. We need to ‘rescale’ this data to get useful results (divide by the sample standard deviation).

# Decathlon Dataset. Step 0: EDA



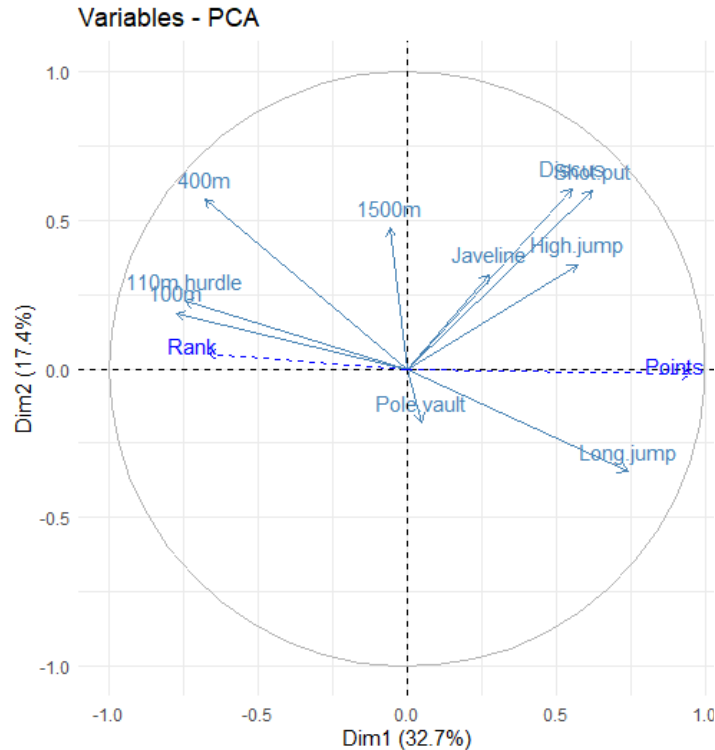
- Clear correlations between the similar field events, and similar track events
- Negative correlations reveal the track events have a different polarity to the field events
  - A bigger number is a good thing in long jump (you jumped further)
  - A bigger number is a bad thing in 100m (you took longer)
  - We decide to multiply all track times by -1 so all events have the same polarity

# A workflow for dimension reduction

0. Identify your variable types and perform appropriate Exploratory Data Analysis
1. Run dimension reduction analysis
2. **Examine the relationships between variables**
3. **Examine the relationships between subjects**
4. Further summarising/interpretation. Choose how many dimensions to keep/examine.
5. Downstream analysis



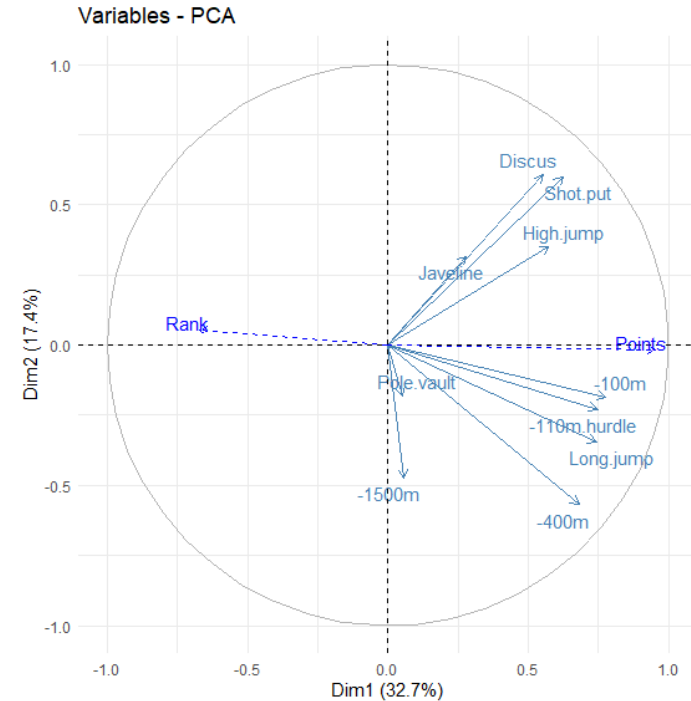
## Decathlon Dataset. PCA step 2: Examine the variable plot



- Let's run the PCA and look at the loading plot
- This loading plot shows what happens if we didn't decide to multiply all track times by -1 (something we picked up in EDA). The 'polarity' problem separates track events and field events on this plot, which makes this plot more difficult to interpret
- Subsequent plots show the solution with the negative track times, and plots relabelled with track events having a minus sign in front

## PCA Step 2: Examine the variable plot

- The loadings plot is where we see the dimension reduction in action. All of the original variables have been projected onto the 2D plane of the two PCs examined in each loading plot (PC1 'Dim1' and PC2 'Dim2' here)
- The correlation of the original variables with each PC is shown by the position of the arrowhead. Drop a perpendicular line from the arrowhead to each axis to see the correlation with that PC (unit circle shows the [-1,1] limits of correlation).
- Those original variables with shorter vectors (e.g. Pole Vault) are less loaded on to the PCs. Recall that the principal directions chosen in PCA preserve as much information as possible, so the variables poorly represented are less important for explaining variability overall in your data

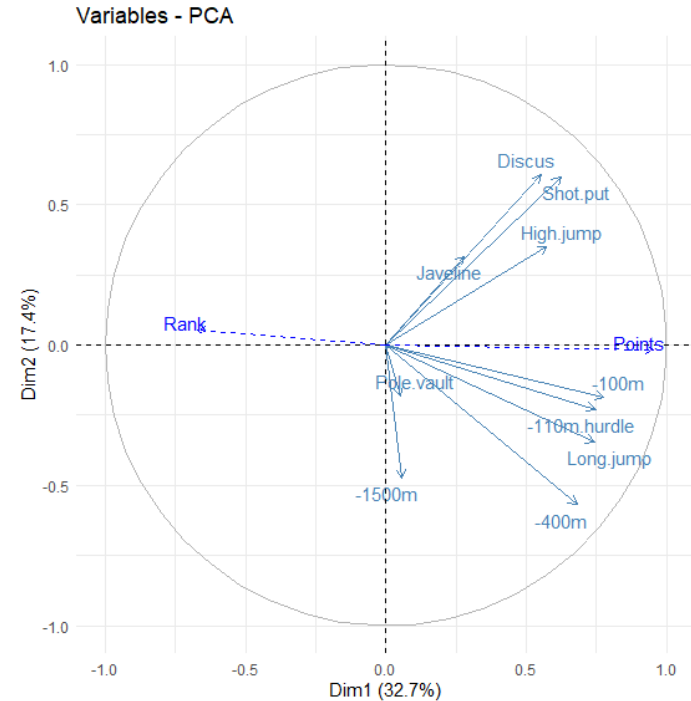


$$PC1 = 0.55(\text{Discus}) + 0.62(\text{Shotput}) + 0.28(\text{Javeline}) + \dots$$

$$PC2 = 0.61(\text{Discus}) + 0.60(\text{Shotput}) + 0.20(\text{Javeline}) + \dots$$

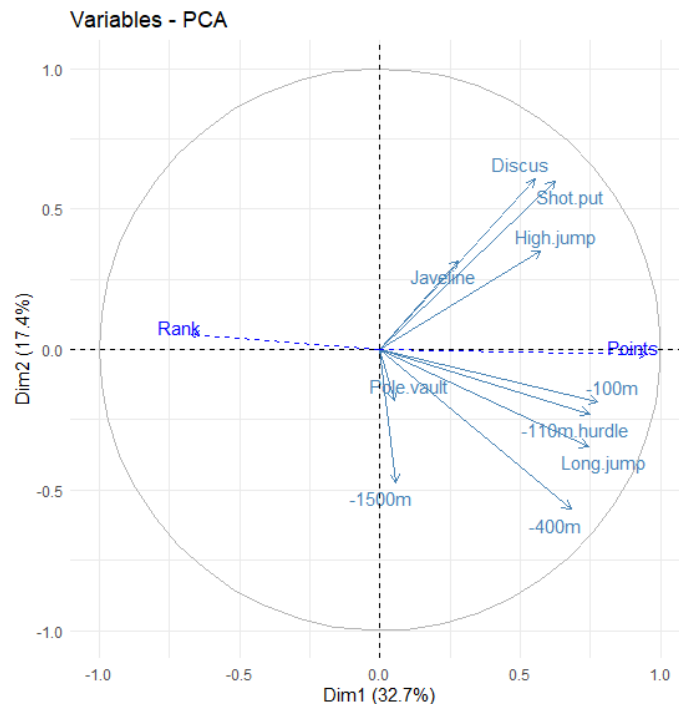
## PCA Step 2: Examine the variable plot

- Sometimes you may have supplementary variables available that are not used in PCA analysis but potentially correlate with PCs. In this case we can use the 'Points' and 'Rank' which for each decathlon give the number of points achieved and the rank of each athlete respectively. These can be added to the loading plot by calculating their Pearson's correlation with each PC.
- Not surprisingly, these are correlated in different directions with respect to PC1. A smaller rank means a better performance, for which more points are awarded.



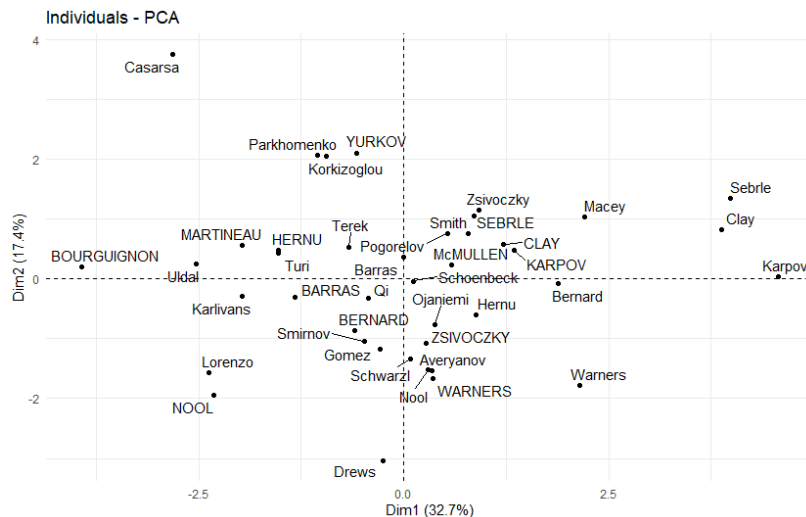
# PCA Step 2: Examine the variable plot

- Interpreting the PCs as latent variables:
  - Clearly PC1 represents performance of the athlete.
    - **A weighted sum** of all events – all variables are positively correlated.
    - Those with higher PC1 scores, have higher points and lower rank.
    - Despite explaining the most variability (33%), this is not particularly interesting as a *latent variable*.
  - PC2 appears to represent specialisation of athletes.
    - It is a **contrast** between events that involve sprinting, and those that don't, which have opposite signed correlations with PC2\*.
    - Sprint specialists achieve by running really fast. Non-sprint specialists achieve most of their points by being good at other field events.
    - Despite explaining about half as much of the variability as PC1 (17%), PC2 has a more interesting interpretation as a latent variable.
    - If specialisation is present in decathletes, this PC will appear reliably in multiple samples.



\*A technical note: the sign of the PC itself is arbitrary! The sign of the correlation between the variable and the PC is not arbitrary.

# PCA Step 3: Examine the subject plot

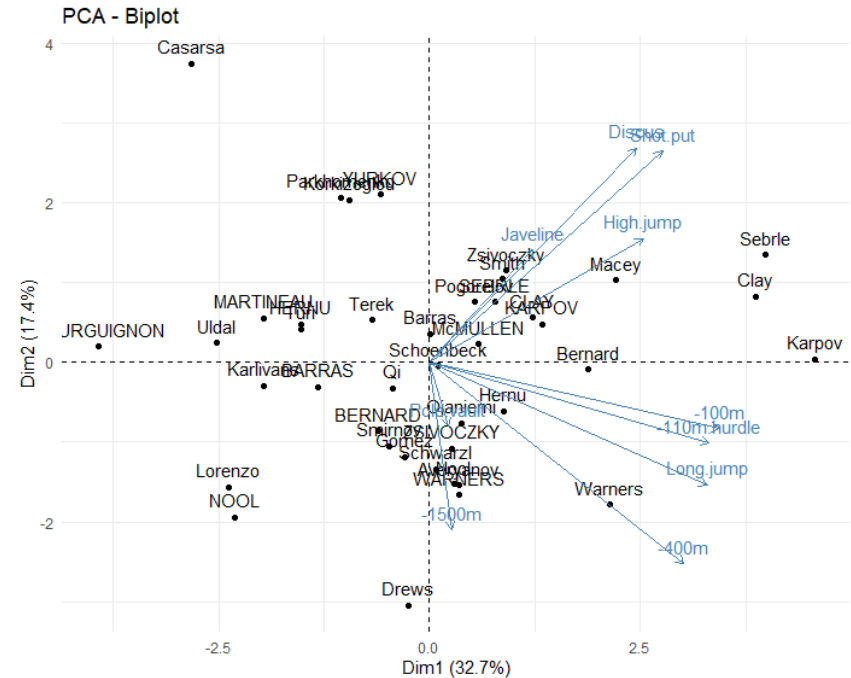


- We can plot the individual observations plotted in PC space, often called the score plot. Their PC scores for two dimensions at a time
- This is again a projection into PC space, but not of the variables, but instead the individual subjects
- Observations close to each other have similar values for the depicted PCs (but not necessarily similar overall). Recall latent variable interpretation for each PC.
- In some examples (not this one) there may be distinct clusters, and you can perform various types of clustering on observations in the PC space.

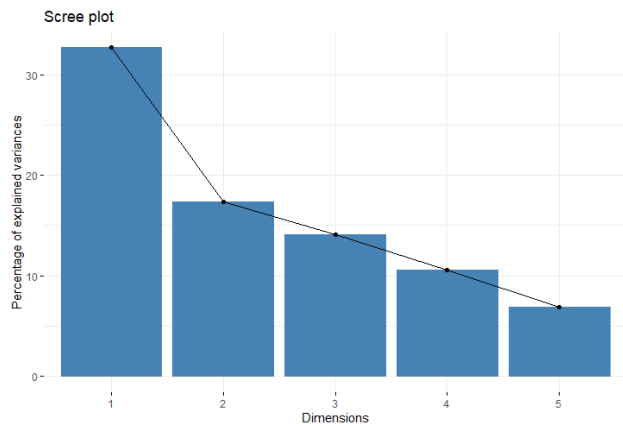


# PCA Step 2 & 3: Examine the variable and subject plot

- The subjects and the variables can be examined on the same plot
- Note that the axes are the same as in the subject plot, and reflect the PC scores
- The original variables can be thought of as additional axes showing the approximate\* + relative position of each of the subjects on the original variable



# How many PCs should we consider?

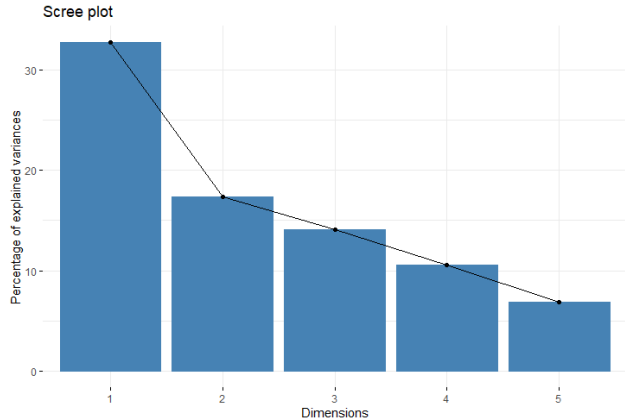


- This ‘Scree plot’ shows the % total variance of each PC. Each PC captures some amount of the total variance (across all original variables = across all dimensions), and this decreases for each subsequent PC
- We have seen that the first 2 PCs capture about 50% of the total variance and we have plausible latent variable interpretation for each. We should definitely consider these two.
- There are as many PCs as there were variables from the output, but we are only looking at the first 5.



# How do you choose how many PCs?

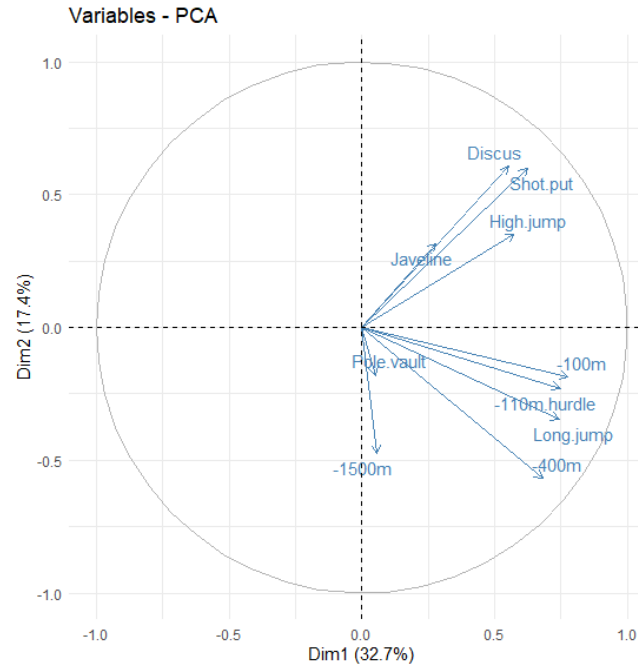
- The steeper the drop between bars, the less information is being captured in the  $k$ th PCs compared to the first  $(k - 1)$  PCs. We can consider dropping the  $k$ th and higher PCs
- Rules of thumb:
  - Smallest # of PCs that together hold 80-90% of variance
  - Keep components with an eigenvalue greater than the average of eigenvalues/Keep components with an eigenvalue  $> 1$  when working with standardised variables as per this example
  - Look for the 'elbow' in the line plot
- If we collected many samples, we would probably find that the first dimensions were more stable in their (relative) directions between samples whereas the last dimensions would keep changing. This would show us that the first dimensions contain more 'signal', and the last dimensions contain more 'noise'.



## Challenge question – correlation between original variable and PC



What is the correlation between -100m (i.e. 100m time with a minus sign) and PC1 score?



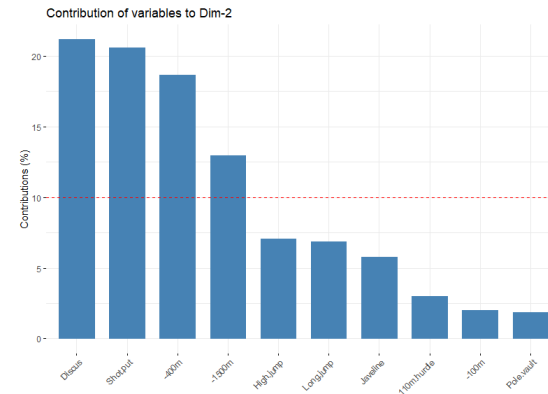
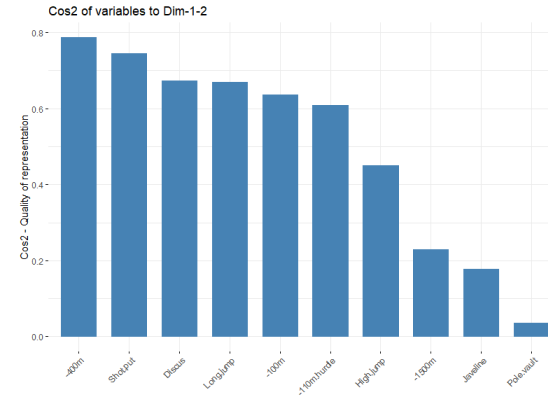
# Challenge question – correlation between PCs



What is the correlation between PC1 and PC2 scores?

# PCA Step 4: Further summarisation

- **Other interpretation aids:**
  - Quality of representation for variables and individuals on a given map (think about the relative length of the vectors on the loading plot)
  - Contribution of variables and individuals to each PC (again the relative length of the vectors, and how important each subject was to the PC output)



# PCA Downstream analysis

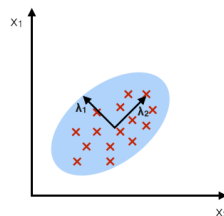
- Possible uses of PCA downstream:
  1. Pre-processing for clustering algorithms
  2. Looking for outliers or time dependency in subjects
  3. As predictors in multiple regression (PCR-Regression)
  4. As a pre-processing step for further dimension reduction

# Method Extension: Linear Discriminant Analysis (LDA)

- PCA can be modified towards a different goal by imposing a different constraint.  
*Recall that in [classic] PCA we want to maximise variability on each principal component*
- If your goal is to find principal components that separate two or more different classes of subject, you could use LDA
- It is a modified version of PCs. LDA calculates PCs that maximise the separability between your classes

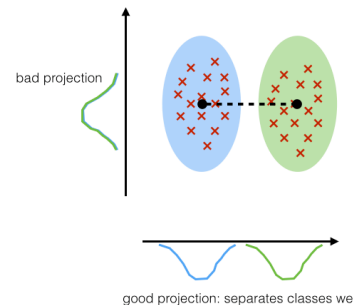
## PCA:

component axes that maximize the variance



## LDA:

maximizing the component axes for class-separation



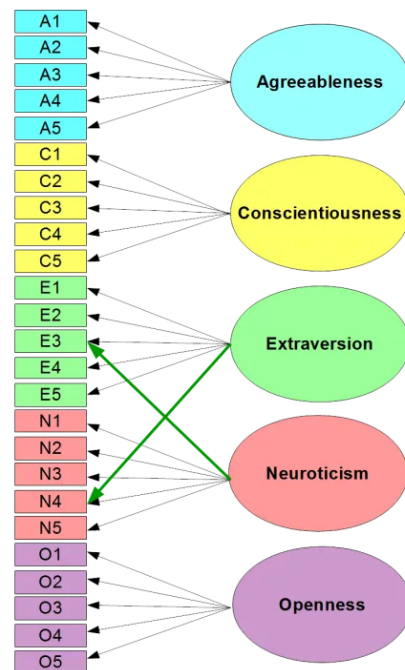
[https://sebastianraschka.com/Articles/2014\\_python\\_lda.html](https://sebastianraschka.com/Articles/2014_python_lda.html)



# Factor Analysis (FA)

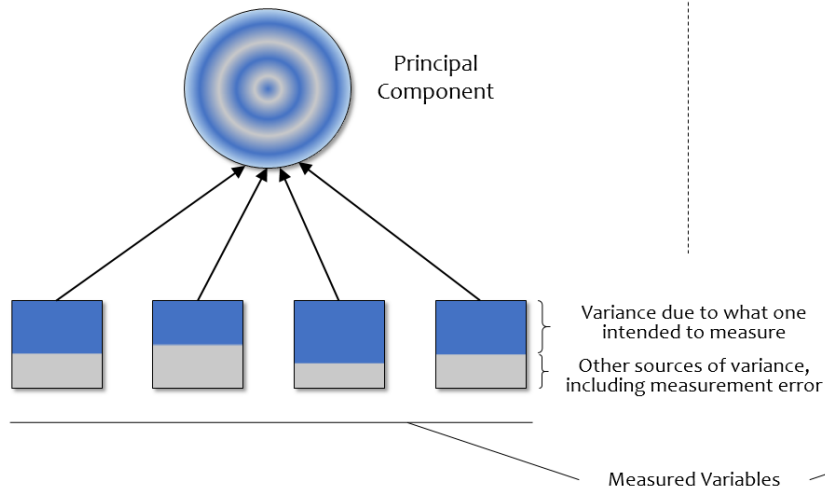
# Factor Analysis

- Another major dimension reduction technique with continuous variables as input
- Similar to PCA in many respects, but differences in the methods reflect different aims and applications
- Most often used to define latent variables with the aim of producing an explicit, and testable model for how these interact with each other and measured variables
- Defining meaningful latent variables is usually more important when using this technique than mere data reduction
- Common in psychometrics. Things like personality traits can't be measured directly but are of interest. Think building a theoretical/functional model using a probabilistic model.

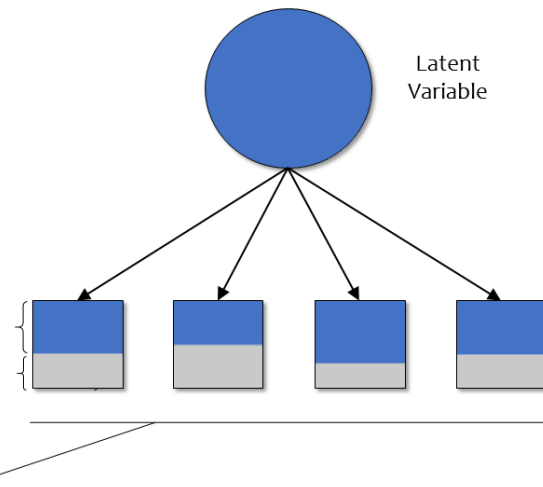


# PCA vs. Factor Analysis

## Principal Components Analysis



## Exploratory Factor Analysis



- PCA identified principal components, which capture patterns of variance in measured variables. You can think of any measurement error as being distributed throughout the PCs (but hopefully relegated mostly to PCs that are discarded)
- In factor analysis you will model error in your measured variables, and this is separate from the shared variance captured by each factor



# PCA vs. Factor Analysis – a terminology note

- In SPSS the Factor Analysis menu lists PCA as a method. But most analysts would not consider PCA to be true Factor Analysis. Factor Analysis would only be when using a ‘factor extraction method’ that partitions variance of a measured variable into unique and common (see previous slide)
- Other differences:
  - Factor analysis uses an iterative algorithm to find an optimal solution, PCA uses direct calculation to find the single solution
  - Factor analysis loading plots do not resemble PCA loading plots, but instead show only the loading for each measured (observed) variable onto each factor
  - Factor analysis usually involves a rotation of the initial solution to maximise interpretability of the factors
- Some people refer to Factor Analysis as Common Factor Analysis, or Exploratory Factor Analysis (EFA) as distinct to CFA (Confirmatory Factor Analysis), or SEM (Structural Equation Modelling) where a defined model is tested for model fit

# Factor Analysis Downstream analysis

- The goal of FA is to come up with an explicit and testable model for how measured variables and factors interact. So FA is usually the first step in a longer process.
- Typical downstream steps from Factor Analysis (FA):
  - Confirmatory Factor Analysis
  - Structural Equation Modelling
- These are referred to in our Surveys 2 workshop, as surveys are commonly used as the input for factor analysis

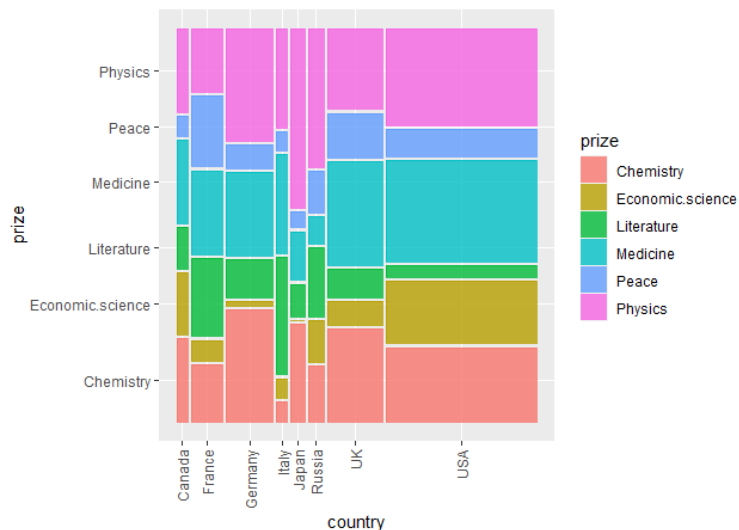
# Correspondence Analysis (CA)

# Correspondence Analysis (CA)

- Very commonly used in surveys with categorical responses.
- Often used as categorical/qualitative analogue of PCA
- Input for PCA and factor analysis is continuous observations
- Input for Correspondence Analysis is categorical observations on two variables: contingency table

# Example: Nobel Prize Data

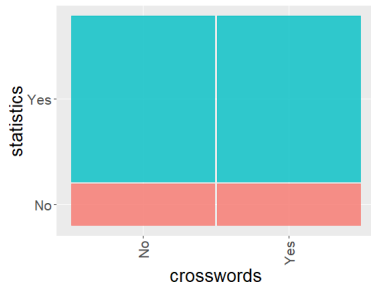
- Nobel prize winners by country from 1901 to 2015. Just G8 countries. Excluding mathematics.
- We could use a contingency table to show frequencies, or a mosaic plot as shown below:
  - The width of the columns represents the column proportions (country)
  - The area of the blocks represents the proportion of all prizes (prize x country)



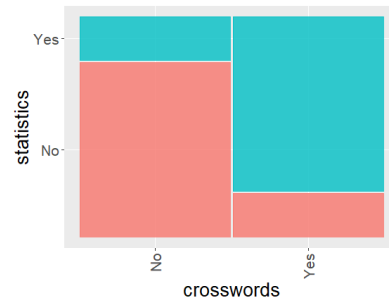


# Association

- In an underlying population, there are two possible relationships between two categorical variables: independence, or association.



INDEPENDENCE  
/no association

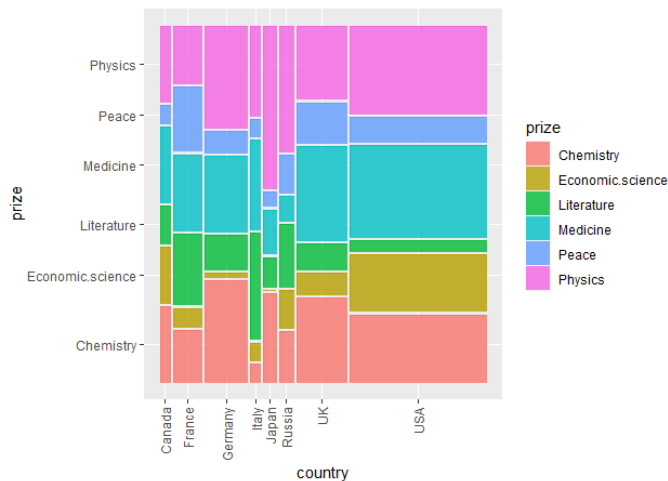


DEPENDENCE  
/association

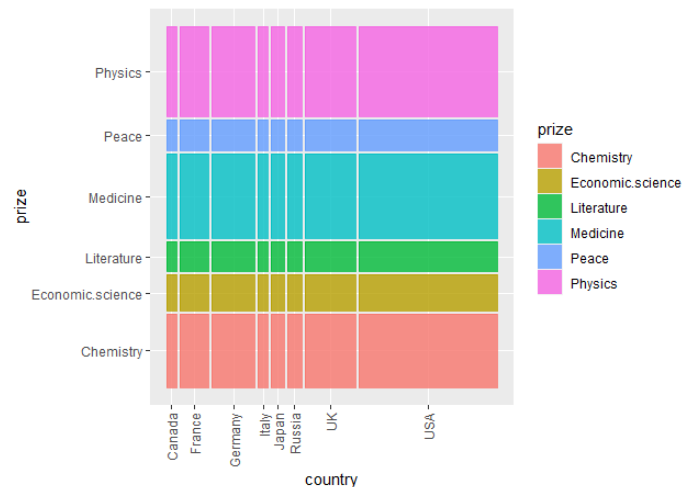
- We use a chi-squared test (or similar) to examine evidence against the null hypothesis of independence for a given sample. This provides a good summary of association for two variables with two categories in each.
- But what about if there are more than two categories in either or both variables?

# Comparing to the independence model

- Is there an association between country and type of Nobel prize won?
- We can compare the observed areas to what is expected under the independence model, where the frequency of each combination of country and prize follows the marginal (overall) frequencies of countries and of prizes



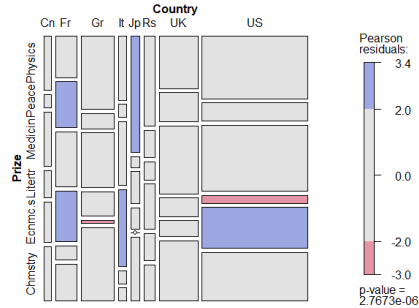
Observed data



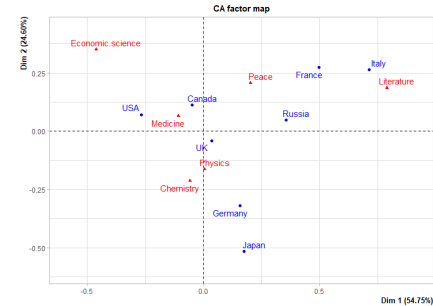
Expected under null hypothesis  
INDEPENDENCE

# Nobel Prize data workflow

- See the R software workflow



Step 0: EDA



Steps 2 & 3:  
'Subjects' and  
'variables' plots

# Method Extension: Multiple Correspondence Analysis (MCA)

- In multiple correspondence analysis we consider more than two categorical variables
- We no longer start with a contingency table of two categorical variable frequencies, but instead with an 'indicator matrix\*' of subjects vs. categories
- Correspondence map can show relation between subjects, variables and categories

|                                |                       |                       |                       |
|--------------------------------|-----------------------|-----------------------|-----------------------|
| How do you take your tea:      |                       |                       |                       |
| Milk                           | Lemon                 | Other                 | Straight up           |
| <input type="radio"/>          | <input type="radio"/> | <input type="radio"/> | <input type="radio"/> |
| Do you have sugar in your tea? |                       |                       |                       |
| Yes                            | No                    |                       |                       |
| <input type="radio"/>          | <input type="radio"/> |                       |                       |



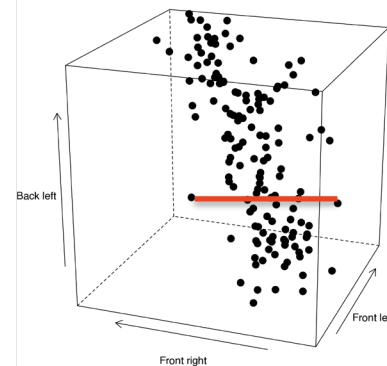
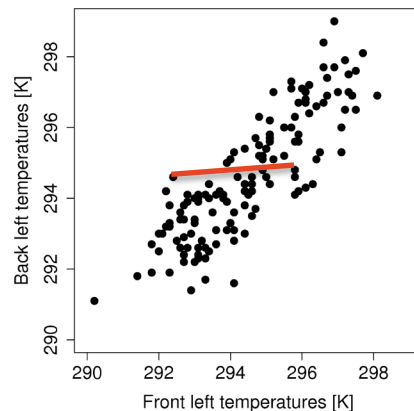
\*Matrix of indicator variables where e.g. (1 = Yes, 0 = No) for the relevant subject x category combination

# **Dimensionality Reduction Techniques:**

## **Distance based**

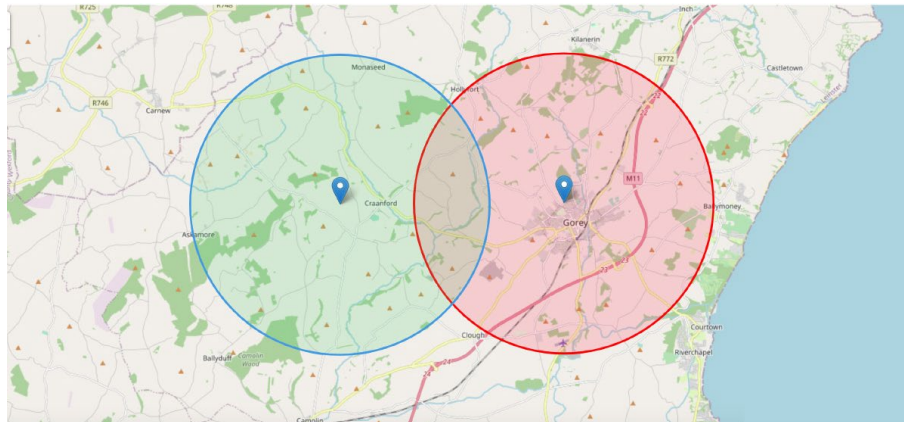
# Distance Based Methods

- The starting point of methods so far have been measurements on **variables**, either continuous (PCA, FA) or categorical (CA)
- Multidimensional Scaling (MDS) and similar methods use the distances between **subjects** as their starting point.
- The goal is to produce a low-dimensional map (often called an ordination or embedding for these methods) that most accurately visualises the distances between the subjects, i.e. preserves as much as possible the distance information from higher dimensions
- Better at accurately capturing local structure (distances between neighbouring subjects), than PCA-like methods, which are better at capturing global structure (distances between all subjects)



# What are distances between subjects?

- We're (very) familiar with geographic distances, but what is the distance between subjects?
- The type of data will dictate the appropriate distance measurement. MDS is very flexible in that it can accommodate different kinds of distances.
  - The Euclidian distance is the straight-line distance between two points in your original variable space
  - Chi-squared distance, is the Euclidian distance between relative frequencies (see the CA example)
  - In ecology, the abundance of species at different sites is measured. Bray-Curtis distance is used, ranges between 0 (identical) and 1 (no similarity)



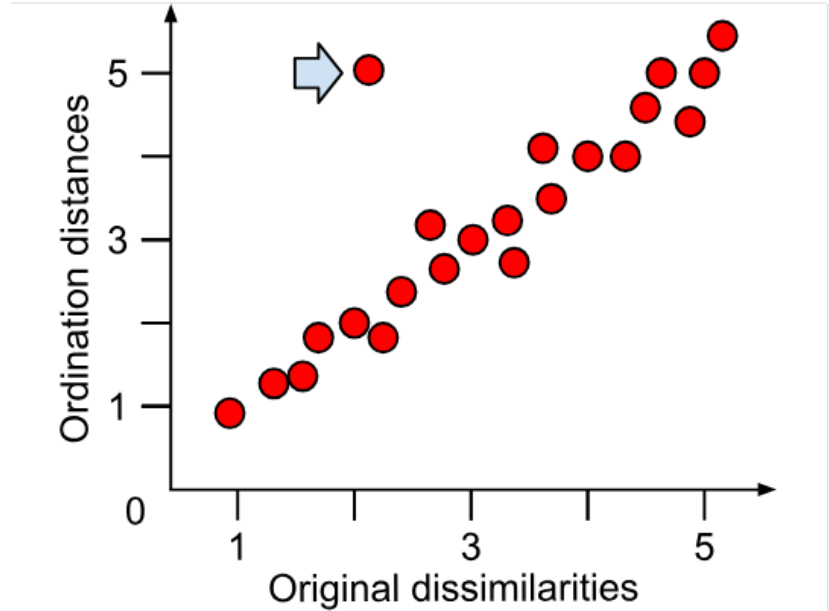
<https://2kmfromhome.com/>

# Multidimensional Scaling (MDS)



# MDS Step 1: Run the MDS

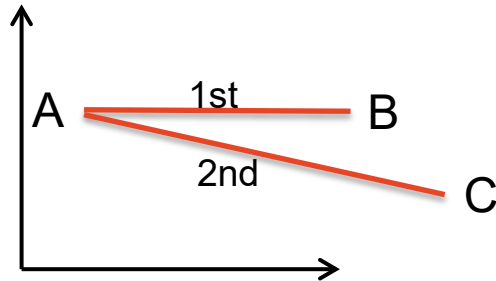
- MDS can be run as either metric or non-metric (nMDS). Metric is also called Principle Coordinates Analysis (PCoA)
- Unlike PCA, it is an iterative algorithm, starting with randomly positioned subjects on the lower dimensional map and stepping towards an optimal solution where the ordination (lower dimensional) distances are as close as possible to the original distances (higher dimensional)



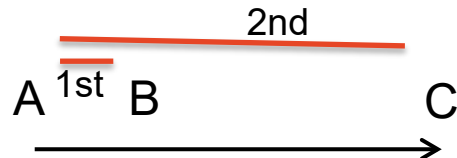
<https://sites.google.com/site/mb3gustame/dissimilarity-based-methods/non-metric-multidimensional-scaling>

# Metric (MDS) or Non-metric? (nMDS)

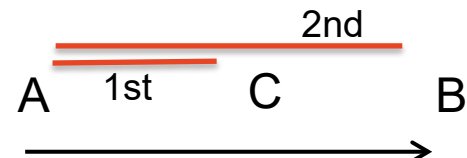
- Metric is appropriate when you expect there can be a linear relationship between the data distances (original variables) and the ordination distances (ordination map). The method and outputs are very similar to PCA. Metric with Euclidian distance *is* PCA.
- In non-metric MDS (nMDS) the success of the ordination depends on preserving the ranks of dissimilarity between subject



ACTUAL DISTANCE  
In higher dimensions



SAME DISSIMILARITY RANK  
In lower dimension



DIFFERENT DISSIMILARITY RANK  
In lower dimension

# Differences of nMDS (and related methods) from PCA

- The output dimensions are not orthogonal to each other as in PCA
- Ordination map is optimised for a pre-specified number of dimensions
  - The algorithm for producing an optimal ordination works by iteration, not calculation of a single solution as PCA does
  - The dimensions are not ordered by variance explained as they are in PCA
- nMDS is able to handle and effectively summarise non-linear relationships



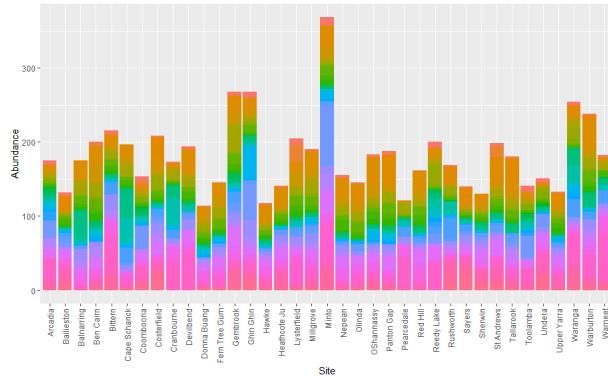
# Example: Mac Nally bird abundance

- Let's take an ecology example for MDS:  
Mac Nally bird abundance data
  - Ralph Mac Nally (1989)
  - Maximum abundance for 102 bird species
  - 37 sites, 5 different forest types (Gippsland manna gum, montane forest, woodland, box-ironbark and river redgum and mixed forest)
- Research question: Do the bird assemblages differ between forest types?

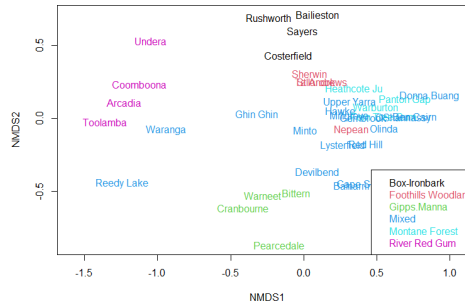


# Mac Nally bird abundance example

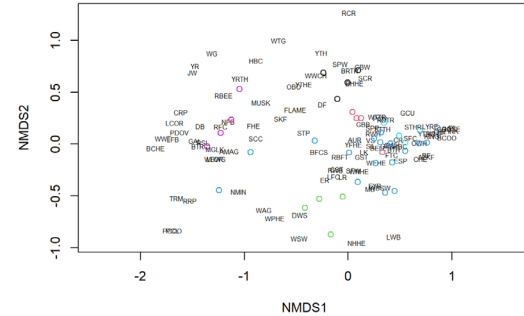
- See the R software workflow



Step 0: EDA



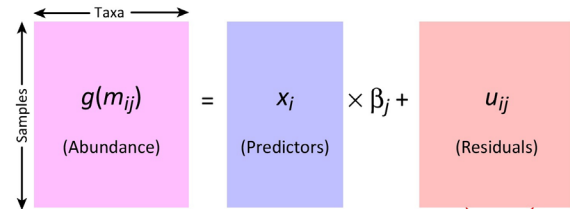
Step 3: 'Subjects' (sites) plot



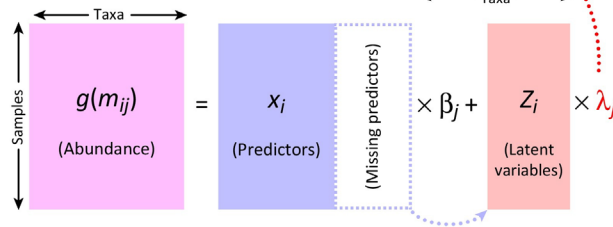
Step 2: Variables (species) plot

# Method extension: Joint species distribution models

(A) Multivariate generalised linear mixed model (GLMM)



(B) Latent variable model (LVM)



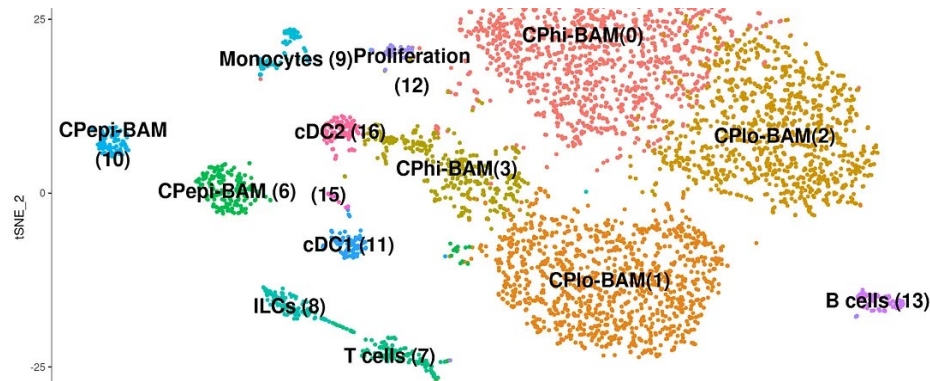
Trends in Ecology & Evolution

- Model-based ordination methods have recently been developed
- These methods handle the challenges of abundance data (non-normality, 0 observations) and allow multivariate [joint] regression models to be fit to this type of data, with environmental variables as predictors
- These methods go beyond dimension reduction techniques, incorporating dimension reduction within a modelling approach



## Method Alternative: t-SNE

- t-SNE is a relatively recently developed method that handles non-linear relationships between subjects
- Often applied to single-cell RNA sequencing
- Tries to make the lower dimensional map match the distance distribution of the subjects in higher-dimensional space



Source: K. Movahedi, Y. Saeys,  
<http://www.brainimmuneatlas.org/tsne-cp-irf8.php>

# Conclusions



# Choosing a method

## EDUCATION

### Ten quick tips for effective dimensionality reduction

Lan Huong Nguyen<sup>1</sup>, Susan Holmes<sup>2\*</sup>

**1** Institute for Mathematical and Computational Engineering, Stanford University, Stanford, California, United States of America, **2** Department of Statistics, Stanford University, Stanford, California, United States of America

\* [susan@stat.stanford.edu](mailto:susan@stat.stanford.edu)

#### Introduction

Dimensionality reduction (DR) is frequently applied during the analysis of high-dimensional data. Both a means of denoising and simplification, it can be beneficial for the majority of modern biological datasets, in which it's not uncommon to have hundreds or even millions of simultaneous measurements collected for a single sample. Because of "the curse of dimensionality," many statistical methods lack power when applied to high-dimensional data. Even if the number of collected data points is large, they remain sparsely submerged in a voluminous high-dimensional space that is practically impossible to explore exhaustively (see chapter 12 [1]). By reducing the dimensionality of the data, you can often alleviate this challenging and troublesome phenomenon. Low-dimensional data representations that remove noise but retain the signal of interest can be instrumental in understanding hidden structures and patterns. Original high-dimensional data often contain measurements on uninformative or redundant variables. DR can be viewed as a method for latent feature extraction. It is also frequently used for data compression, exploration, and visualization. Although many DR techniques have been developed and implemented in standard data analytic pipelines, they are easy to misuse, and their results are often misinterpreted in practice. This article presents a set of useful guidelines for practitioners specifying how to correctly perform DR, interpret its output, and communicate results. Note that this is not a review article, and we recommend some important reviews in the references.

#### Tip 1: Choose an appropriate method



#### OPEN ACCESS

**Citation:** Nguyen LH, Holmes S (2019) Ten quick tips for effective dimensionality reduction. PLoS Comput Biol 15(6): e1006907. <https://doi.org/10.1371/journal.pcbi.1006907>

**Editor:** Francis Ouellette, University of Toronto, CANADA

**Published:** June 20, 2019

**Copyright:** © 2019 Nguyen, Holmes. This is an open access article distributed under the terms of the [Creative Commons Attribution License](https://creativecommons.org/licenses/by/4.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

- I've given you a broad overview of basic dimension reduction methods
- There are many extensions to these methods and many other methods available
- Many excellent tutorials and other resources are available for implementing a particular method
- Also consider the choice of method after doing some basic EDA

Nguyen LH, Holmes S (2019) Ten quick tips for effective dimensionality reduction. PLoS Computational Biology 15(6): e1006907. <https://doi.org/10.1371/journal.pcbi.1006907>

# Further Assistance at Sydney University

## SIH

- [Statistical Consulting website](#): containing our workshop slides and our favourite external resources (including links for learning R and SPSS)
- [Hacky Hour](#) an informal monthly meetup for getting help with coding or using statistics software
- 1on1 Consults can be requested [on our website](#) (click on the big red 'contact us' link)

## SIH Workshops

- Create your own custom programmes tailored to your research needs by attending more of our Statistical Consulting workshops. Look for the statistics workshops on [our training page](#).
- [Other SIH workshops](#) including machine learning workshops for hands-on DR
- [Sign up to our mailing list](#) to be notified of upcoming training

## Other

- Open Learning Environment (OLE) courses
- [Linkedin Learning](#)

# How to use our workshops

Workshops developed by the Statistical Consulting Team within the Sydney Informatics Hub form an integrated modular framework. Researchers are encouraged to choose modules to **create custom programmes tailored to their specific needs**. This is achieved through:

- **Short 90 minute workshops**, acknowledging researchers rarely have time for long multi day workshops.
- Providing **statistical workflows applicable in any software**, that give **practical step by step instructions which researchers return to when analysing and interpreting their data or designing their study** e.g. workflows for designing studies for strong causal inference, model diagnostics, interpretation and presentation of results.
- Each one focusing on a specific statistical method while also integrating and referencing the others to give a **holistic understanding of how data can be transformed into knowledge from a statistical perspective** from hypothesis generation to publication.

For other workshops that fit into this integrated framework refer to our training link page under statistics <https://www.sydney.edu.au/research/facilities/sydney-informatics-hub/workshops-and-training.html#stats>

# We recommend our Experimental Design and Sample Size Workshops

## Experimental Design Workshop

- Far too many researchers think they know all they need to in this area. We commonly see designs that could be substantially improved for stronger causal inference and improved results which leads to publication in higher impact journals (amongst other benefits).
- Even if you have already collected your data it is well worth attending since it may improve your write up and analysis e.g. we had a client who didn't realise they had a very strong Before/After Control/Impact (BACI) design.

## Sample and Power Workshop

- Shows the steps and decisions researchers need to make when designing an experiments to ensure sufficient sample e.g. Power, minimum required to fit the necessary model, etc.
- Also how much Power the study has i.e. does it have sufficient power to detect the effects you expect to see, or is your study a complete waste of time and resources.



# References

- Tutorial examples used:
  - [PCA Decathlon example](#) from [FactomineR](#)
  - [CA Nobel Prize example](#) from [François Husson's github](#)
  - [nMDS Mac Nally example](#) from Murray's R resources  
[Currently offline]
- **See the software workflow for this workshop for R code and tutorial data with these examples**

## A reminder: Acknowledging SIH



All University of Sydney resources are available to Sydney researchers **free of charge**. The use of the SIH services including the Artemis HPC and associated support and training warrants acknowledgement in any publications, conference proceedings or posters describing work facilitated by these services.

*The continued acknowledgment of the use of SIH facilities ensures the sustainability of our services.*

### **Suggested wording for use of workshops and workflows:**

*“The authors acknowledge the Statistical workshops and workflows provided by the Sydney Informatics Hub, a Core Research Facility of the University of Sydney.”*

# We value your feedback



We want to hear about you and whether this workshop has helped you in your research. What **worked** and what **didn't work**.

*We actively use the feedback to improve our workshops.*

Completing this survey really does help us and we would appreciate your help! It only takes a few minutes to complete (*promise!*)

You will receive a link to the anonymous survey by email



# Glossary of Terms

- **Variance** is simply variability of a single variable. Similar to modelling, unexplained variation is a bad thing because it contributes to the error of our estimates, but explained variation is a good thing because we can attribute that variability to a particular factor, and turn ‘noise’ into a meaningful ‘signal’. When we try to “maximise variability” in PCA, you can think about this as trying to maximise explained variability or “information”.
- **Total Variance/Total inertia** usually refers to the variance or inertia summed across all variables/dimensions
- **Dimension/Components/Factors** Dimension reduction techniques employ linear algebra, which often uses terminology from geometry. Mathematically, we can think about each original variable, or synthetic variable as being a dimension (I use ‘dimension’ to refer exclusively to the dimensions identified by the various dimension reduction methods These dimensions have differing names depending on the technique being used.)