

多次元尺度構成法 ——その基礎と広がり——

小杉考司（専修大学）

はじめに

今日は多次元尺度構成法, Multi-Dimensional Scaling (以下 MDS) のお話をします。今日のお話ですけれども, 前半の方で基本的な MDS のお話をした上で後半の方で特殊なというか, より進んだ MDS のモデルについてお話しさせていただきたいと思っています。

MDS の基礎

MDS の特徴, 長所, 短所

ということでまず MDS って何なのだという話と, 改めて距離って何なのかみたいところの話を進めていきたいと思います。一言で MDS の特徴は何かと言うと, 距離の情報から地図を作ることだと思ってください。距離の情報から地図を作る, 言ってしまうとこれだけの話です。これがなぜ良いのかという話ですけれども, 大きな特徴の一つとしてはですね, 被験者への負担, 介入過程が少ないデータからスケールを作ることができることです。いろんな心理的な刺激を出している人々に評価してもらってそこからデータを作っていくということを, 心理学ではやっているわけですが, その元になるデータが距離, 似ているか似ていないか, という判断をさせるだけで分析ができる。これが MDS の醍醐味です。もうひとつは, その似ている・似ていないという程度の判断が, 順序尺度水準でもいいということです。心理学のデータのほとんどは, 例えば, 因子分析法だとか構造方程式モデリング (SEM, あるいは共分散構造分析) のようなデータなんかは間隔尺度水準以上の量的なデータであることを仮定しますが, その仮定が必ずしも成り立っていないというようなときであっても, 数値を与えることができるというのが面白いポイントですね。三つ目に, 結果の空間に加筆修正できる自由度の高さがありますが, この点については今日のお話の後半部分の応用編のところでお話できれば, というように思います。

また, 先に欠点の方を言っておきますと, データにもよりますが, 結果が多少不安定なことが挙げられます。距離というのはその条件に合うものであればなんでもいいのですが, そのラフさ故に, 距離をどう計算するか, その計算方法が違えば結果も変わってしまうということがあるんですね。誰がどこからどう見ても間違いない本当の真実, というようなものを期待されると, ちょっと「思ってたんと違う!」というようなことになるかもしれません。

いきなりですが, ちょっと実演してみましょう。統計ソフト R を使って今日はお話しさせていただきます。その R のサンプルデータの中に **eurodist** というデータセットがあります。これはヨーロッパの都市の間の距離がデータになっているもので, データを見ていただいたらですね, 例えばアテネとバルセロナとの距離が 3313, なんて書いてありますね。単位はキロメートルです。行と列がずらっと並んでいるので, 画面の下の方にも出力が続いているのですが, 距離行列というのは基本的に対称行列です。つまりアテネとバルセロナの距離はバルセロナとアテネの距離といっしょになりますので, 表示するときは下三角行列, つまり全体の距離行列の下半分だけ表示するようになっています。これをいきなりですが, **cmdscale** という, MDS の関数に与えます。これも R に最初から入っている関数で, MDS を実行しろ, つまり地図を作れというコードを走らせると, こんな結果が出てきます (Figure 1)。

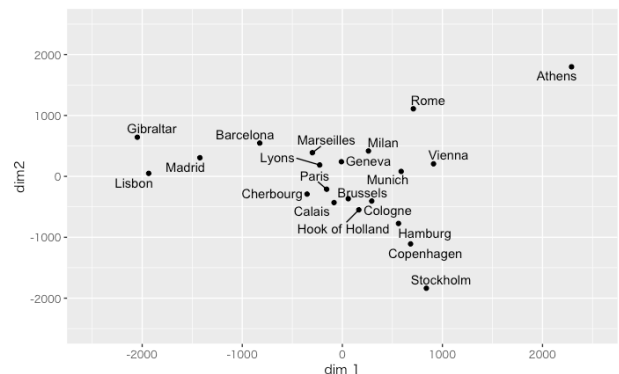


Figure 1. MDS による出力例 (ヨーロッパの都市間距離より)

これを見ると右上の方にアテネがあって, 下の方にストックホルムが出てきてます。確かストックホルムって北の街じゃなかったかな, と思う方もいらっしゃると思うんですが, このように感じで相対的な位置関係が出るだけですので, 必ずしも上を北にした結果というのが出るわけではありません。そこで手動で, これをぐるんと反転させてみます。そうすると上にストックホルムがきて右下にアテネが来て, アテネのちょっと離れたところにローマがあつたりして, こういう, よりヨーロッパの地図に似た図が出てくるかと思います (Figure 2)。

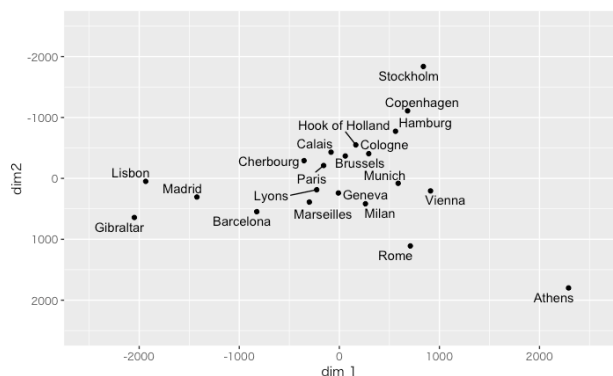


Figure 2. MDS による出力例 (第二軸の反転)

こちらが Google Map を使って「ヨーロッパ」で検索したときに出てくる画面です。多少倍率は変えましたが、これと先ほどの結果を重ねてみますとこんな感じになるわけです(Figure 3)。これは私がゴソゴソと、角度をちょっとずつ変えて透過率を上げながらやって、重なるように描いたものなんです。例えばアテネとストックホルムの辺り、実際の地図が青いポイントです。MDS の重ねた地図が赤いポイントなんですが、その辺が重なるように角度をつけたりしてですね、うまく合わせると、ストックホルム、コペンハーゲン、アテネ、ローマ、ジェノバ、マドリッド…はちょっと外れてますけれども、とにかく、ちょっと調整すると、なるほど地図が書けましたね！

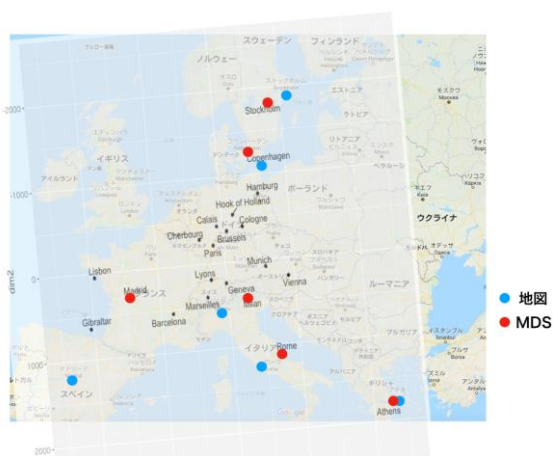


Figure 3. MDS の結果と実際の地図を重ねてみた

これを見て完全に一致というのは言い過ぎかもしれませんが、元々がそれぞれの都市間がどれぐらい離れているかという距離の情報だけでして、そこから地図を作っていますので、そうだとするとそこそこの再現ができてるといっても良いのではないのでしょうか。完全再現というには無理があるとしても、そこそこできてるんじゃないかと思っていただければ、MDS に少し興味を持っていただけるんじゃないかと思います。本当は、私もこれ、もうちょっと精度が上がるものかな、と思ったんですけれ

ども…。例えば Google の地図もですね、最近なんか補正が入って距離が正確になるように、地球の曲面の率に合わせたような地図の出方になったそうです。MDS の方の図は、これはデフォルトでは二次元マップとして出力するのですが、実際は三次元の世界観を二次元のマップに落としてるわけですから、なかなかぴったり完全に再現、というわけにはいかないというところがあるんです。そのあたりを勘案してもらって、まあそこそこ再現できているのではないかというふうに思っていただけなのです。と思います。

もし、「なるほど距離から地図を再現するのか」、という筋道を納得していただきましたら、次にやっと本題の、じゃあ距離って何なのか、あるいはこの MDS の関数の中でどうやって作っているのかという話になりますね。そこで少し遠回りに思うかもしれませんが、距離って何かというお話を今からしていきたいと思います。

距離ってなんだろう

私の好きな『伝染んです。』(吉田,1994)という漫画に大好きなシーンがありまして。四コマ漫画なんですけどね、一コマ目で「おっ、距離！」って喜んでいるおじさんが出てくるんです。それを周りのみんなで、「いい距離だ」「ほう、いい距離ですね」「私ら距離マニアを満足させる距離はめったにありませんからなあ」なんて言ってるわけです。そこへ最初の人の奥さんらしき人が、「あなた！ そんなわけのわからない距離道楽はもうやめてください！ 店のことも子供のこともみな私に押し付けて。うう……」って出てきて、「す、すまん」とかって謝るんですけど、その時もいい距離を保っているという(笑)。まあ不条理漫画なので意味はないんですが、距離というこの漫画がいつも思いだされるんです。皆さんもこのマニアたちに負けず劣らずですね、距離というものにもっと興味を持っていただければな、と、そんな風に思っています。冗談です。

さて、距離って何なのかという、数学的なところに戻って考えますと、ルールはそんなに多くありません。一つ目は2点 x と y の距離を $d(x,y)$ と表すと、距離は必ず正の値を取りますよ、っていうことです。二つ目は x と y の距離は y と x の距離と等しいですよ、これが対称性の条件というやつですね。で、三つ目は、もし x と y が同じ、つまり同じ点であればその距離はゼロということです。アテネとアテネの距離は 0 km みたいな意味ですね。で、最後の四つ目は、 x と z 、 z と y という x,y,z の三つのポイントがあったとすると、 x と y の距離は直線の距離はどこか一箇所、 z を経由したときの距離に比べて必ず短くなるというやつです。これらのルールが満たされていけば全て距離と考えることができるというのが距離の一般的な公理です(Figure 4)。

距離の公理

- 2点x とyの距離を $d(x,y)$ とすると,

$$d(x, y) \geq 0 \quad \text{非負性 (正定値性)}$$

$$d(x, y) = d(y, x) \quad \text{対称性}$$

$$x = y \Rightarrow d(x, y) = 0$$

$$d(x, z) + d(z, y) \geq d(x, y) \quad \text{三角不等式}$$

- の条件を満たせば全て「距離」。

Figure 4. 距離についてのスライドより

この条件を満たす様々な距離が考えられているんですが、実はRには距離に関してはdist関数というのがあります。最初から距離を計算してくれる関数が用意されています。この関数を使うとデータ行列を与えると答えとして距離行列が返ってくるので非常に便利な関数ですね。ちょっと例として、iris というデータで関数の使い方をお示ししますね。この iris というのもRに入っているサンプルデータですが、4つの変数に関して150個体のデータがあります。これの距離を計算しなさい、とdist関数に与えますと、下三角の行列が返ってきます(Figure 5)。



Figure 5. 統計環境Rによる距離の関数

この行列はサイズ150の下三角行列ですが、このデータを転置してから関数に渡すとサイズ4の下三角行列になります。とにかく三角行列として二つの変数あるいはそれぞれの個体間の距離が返ってくるのがdistという関数です。このdist関数には様々なオプションがありまして、特に何の指定もしなければ、デフォルトではユークリッド距離が計算されます。その他にもマキシмум距離だとか、マンハッタン距離、キャンベラ距離、バイナリー距離、ミンコフスキー距離という6種類もの距離が用意されています。なので、皆さんもどんな距離使おうかしらと思

いながら、興奮しながら聞いていただきたいと思います。

もちろんどのオプションを選んでもらっても、それぞれ距離の公理が満たされるよう数字になっています。最初のデフォルトのユークリッド距離というのは、皆さんもよくご存知のいわゆる距離、一般的な距離です。二次元であればピタゴラスの定理で表わされるように、直角三角形の斜辺以外の二つの辺を二乗して足し合わせ、平方根が斜辺になる、という形で表現されるのがユークリッド距離ですね。二次元の場合はこうですが、三次元、四次元、五次元の場合であっても、各次元について二乗した辺を次々足していきます。一次元目を二乗、二次元目を二乗、三次元目を二乗というように、全部二乗して足したものの平方根を取ったらユークリッド距離になります。これがまあいわゆる距離の一番わかりやすい例です。

二つ目にマキシмум距離というのがあります。これは要素同士の差の絶対値が最大のものを距離とする、いろんな要素の中の絶対値が最大になるものを距離にしましょうということで、マキシмум距離という名前がついています。

マンハッタン距離というのも結構有名ですね。直線ではなくてブロック化された都市の距離みたいな、まっすぐ行かない距離、絶対値の和で、xとyのそれぞれの次元の絶対値の和で計算するというのがマンハッタン距離です。次にこのキャンベラ距離ですね。これ何でキャンベラっていう名前なのかかわかんないんですけど、マンハッタン距離の拡張みたいな形で、定義としては足したやつと絶対値分の引いたやつと絶対値の和みたいな形で表現されるものです。

あるいはバイナリー距離というやつで、これは0・1のビットデータなんかに関する距離の計算の仕方、一方が1のときに他方も1であれば近い、一方が1のときに他方が0であれば遠い。一方が0のときはノーカウントにするというルールで計算される距離です。

最後に、ミンコフスキー距離です。先ほどいろいろお話してました距離の計算をするときってというのは、例えばユークリッド距離は二乗したものを足して平方根をとる、としましたよね。この累乗の計算を、p乗したやつとp乗根(1/p乗)を取るというように一般化すれば、それがミンコフスキー距離と言われる距離になります。二乗するとか、三乗、四乗、五乗、好きな数字を入れていただいたミンコフスキー距離というのを作ることができます。特に三乗とか四乗とかっていうのを目標にして距離にするというのは僕もあんまり聞いたことはないんですが、距離の概念はこういう風にして何乗するかという形で一般化できますよという話があるわけです。ちなみにこの考え方なんですけど、一乗をしたときは距離の長さ、二乗したときは円の状態、無限大にすると一番端っこだまで行きますよということでこの一番端っこの優勢次元距離というのが実は先ほどのマキシмум距離と同じになっているとすることができます(Figure 6)。

method="minkowski"

・ 一般化された距離

・ $p=1$ ならマンハッタン

・ $p=2$ ならユークリッド

・ p はオプションで指定可能(デフォルトは2)



$$d(x, y) = \left(\sum_{i=1}^n |x_i - y_i|^p \right)^{\frac{1}{p}}$$

Figure 6. ミンコフスキー距離に関するスライドより

こういう風にして距離と言ってもいろんな形で取ることができるんだということを分かってください。同じデータであってもそれがユークリッド距離で計算した場合とマンハッタン距離で計算した場合とでは距離の計算が変わってますから、当然 MDS の結果にも影響されるわけです。もっと言うと相関係数とか共分散というのも、相関係数は一方が高ければ他方も高い、一方が低ければ他方も低いというような意味になりますので、これも類似度を表す指標になっているんじゃないかというように考えることもできます。例えば、相関係数はマックスが1になってますので、1から相関係数の絶対値を引き算したらそれも非類似度。つまり距離を表してると考えることもできるわけです。正の相関の場合ですけど。ともかく、こういうようなものを非類似度、距離のデータとして分析することも可能なんです。

こういう意味で MDS の計算のもとになる距離っていうやつがいろんな定義の仕方があり、多様なものを距離とみなすことができる、みなして分析できるというのが MDS の強さでもあるわけです。

距離に関する心理学のデータ

心理学に限ってこの MDS の強みというのを考えますと、類似性の判断、あるいは直接、非類似性を判断させるというような課題を課してもいいですが、類似性の判断をしたものを距離とみなして分析に使うことが一般的です。

直接何か測るのではなくて、類似性というのをデータにすること、どれだけ利点があるか、と思われるかもしれませんが、いくつかあるんです。まず一つ目は、似ているとか似ていないという包括的な判断をさせるだけで良いということ。何かの対象について、例えば「ミネラルウォーター」について、これとですね「綾鷹」みたいなのがあったとして、似ていますか・似ていませんかって聞いたら、誰だってなんとなくうん似てるとか、いや水とお茶だから違うとかっていう判断はしやすいわけ

ですね。心理学者が尺度を作って、水の透明度はどうですかとか、味の深さはどうですかとか、飲んだ後の感触はどうですかみたいな項目を考え出してですよ、それを一つひとつ評価してもらって、その差分を計算して類似度を算出するというのもいいんですが、それは言い方を変えると実験者とか調査する側がもう評価次元を決め打ちしちゃってることになります。本当は、被験者はそんな次元で、細かい透明度とか、味の深さとか、のどごしとかの次元で判断してるんじゃないかと、何となく似てるとか何となく似てないというような、漠としたイメージしかないのだけど、それを下位次元に落とし込んでいくとひょっとしたら大事な何か壊れたりとか、なくなってしまうかもしれない。言語化できない感覚、というのがありますから、そういうときに「似ていますか、似ていませんか」というだけで聞けるというのがメリットの一つになるのではないかと思います。被験者への認知的負担も少ないです。

そのことと関連するのですが、全体的に似ていますか・似ていませんか、非類似度どれぐらいですか、という感じで聞きますので、社会的望ましさ、つまり、あまり悪い言葉は使わないでおこうとか、ネガティブな表現は避けようとかっていうような、調査データなんかの場合はそういう望ましさ判断みたいなのが入ってくると思いますが、それを取り去った自由なデータを入力することができる。この辺がメリットの一つではないかと個人的に思っています。ここで、例えばこのお茶とこのお茶は似ていますか・似ていませんか、何となく似てる・何となく似てないというような感じで答えさせるデータを積み重ねていったとすると、それはちゃんと考えてないデータだからまともなデータなの？みたいなツッコミがあるところもあるかと思います。もし厳格な手続きを取ってやるのであれば、複数回同じ刺激対を出して、そのときにちゃんとその回答に一貫性があるかどうかというのを確認しながら進めていくというやり方もありますね。最初は「おーいお茶」と「綾鷹」はまったく似てないと答えていたのに、「綾鷹」と「おーいお茶」という組み合わせでやったらとても似ていると答えてた、回答に一貫性がなかった、というようなことがあったら、何巡かしてからもう1回提示して、収束するとか答えが一致するまで答えさせるっていうようなやり方をしたりします。そもそも、判断するときは、おそらく被験者の中では言葉にできないけど、何か一貫した評価次元で似ている・似ていないというのがあるはずなので、それを十分にチェックしておけば類似性のデータであっても十分に信憑性のあるデータが手に入ると考えられます。

この辺の話はですね、高根先生の『多次元尺度法』という東京大学出版会が出してる本(高根,1980)に詳しいです。おそらく大分古い本なのでもう絶版になってるんじゃないかと思うのですが、データの取り方なんかについてはこの本に非常に詳しく書かれていますから、興味がある人は図書館とか古本屋さんとかで是非この高根先生の『多次元尺度法』という本を探してみてください。

ださい。古典的な技術について非常に丁寧に書いてありますのでとても勉強になる本です。

評定尺度法 心理的なデータとしての距離ですが、先ほども言いましたように、評定して取るということ、つまり類似度評定で「似ていますか、似ていませんか」という形で聞くことももちろんできます。例えば、「日本人とユダヤ人」と言う2つの刺激を出して、これが似ていますか・似ていませんかっていうのを25段階の目盛りで聞くこともできます。尺度の目盛りをつけるなら、左端が非常によく似ているで、右端がまったく似ていない、というようにするんです。どの辺ですかというように答えてもらってもいいですし、25段階って大変だと思うのであれば、いわゆる5段階、7段階ぐらいの、「非常に似ている」から「やや似ている」といったカテゴリをつけて、段階を落としていただいて。そこまでいくと、ひょっとしたら順序尺度水準のデータになっているかもしれませんが、とにかくそのように取る方法もありますね。他にも、質問紙の紙に直線を引いておいて、両端だけ提示しておいて、該当する程度の線上にチャットとチェックさせて、その長さを測って、何センチぐらいの所にチェックをしたかということで類似度とするというような使い方もできます。最近ではWeb調査なんかでもできますし、その中で答えさせ方としてスライダーみたいなものがありますから、スライダーでどの辺で似ているとか似てないと回答者に調整してもらって、距離データとみなすというような使い方も可能かと思います。

心理指標を類似度とみなす場合 こういう評定尺度法ももちろんあるんですが、その他に、刺激の混同率なんか也使えます。昔、シェパードさんという人がやったデータの例ですが、モールス信号の聞き間違いの多さみたいなのをデータにするというのがあります。例えばaっていうのがトントンツでbっていうのがツートントンだったとすると、それを聞き取るときにaをbと聞き間違えた回数がすごく多いのであれば刺激が似ているわけです。似ているから聞き間違えたのだ、なので類似度が近いと考えます。あるいは、もうこれとこれとはまったく取り間違えようがないですよというような刺激であれば、これはもう似てないのだから距離が遠いというようにして、混同率をもとに距離のデータにすることもできます。

あるいは記憶の実験の単語なんかの話によくあるように、この値がどれくらい代表的なのか、を距離にすることもあります。例えばりんごっていうのは果物の中でとても代表的で、例えばそうじゃない名前、代表的でない果物がぱっと思いつかないですけど、そうですね、ドリアンとかがあったとしたら、これはりんごに比べてだいぶ遠いと考えます。また、これと連想価が高いものは似てるけれども、あんまり連想で出てこないものは遠いとして考えるというような、そういうものを距離と考えることもできます。

あるいは刺激の般化勾配ですね。Sという刺激があればRという反応をするという条件づけがある中で、Sという刺激に対し

てPという反応をしちやった、というこの間違える確率みたいなのを類似度のヒントにするわけです。間違えやすいということとは刺激を取り間違えやすいということだから、刺激の距離が近いと考えるとか。同じか違うかというのを判断させる課題において、「同じ」とする判断が早ければ類似性が高い、「違う」とする判断が遅ければ類似性が近い、といったようなそういったものを距離と考えることもできます。

あるいは社会心理学的な、この辺から僕はMDSの道に入ったのですが、友人関係ですね、学級の中で誰とよく遊ぶの？三人ぐらい名前を挙げてごらんとかっていう、その誰と誰が仲良しですかというような人間関係のデータ、社会ネットワークのデータみたいなやつを距離だとみなして分析するというようなこともできます。

こういう風に考えていくと、距離の公理が満たされてそうな、距離と見なせそうな心理学的な刺激は、結構いろいろあるのではないのでしょうか。しかもその背後にあるかもしれない、下位次元一つひとつの細かな判断まで求めずに、被験者に漠と「似ていますか・似ていませんか」とか、イエスカノーカ、同じか違うか判断してくださいというような形で取ることもできたりするので、そういう意味でデータの間口は非常に広いんですね。この辺の何をデータとするかというのは、実験者の工夫にかかっているところですので、ぜひいろんな工夫をしてですね、どんどんMDSを使っていただければ嬉しいです。

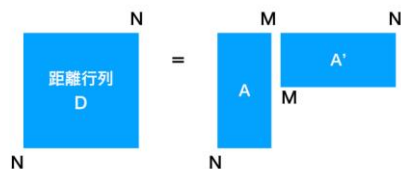
数理的なカラクリ

さて、これで何を距離とするかという話がわかりましたので、どうやってこの座標を作ってるんだ、という具体的なお話に進んでいきたいと思います。最初にYH定理って書いています。これ、笑い話なんですけど、私、前いた大学で英語の本を読んで翻訳させるという授業をやったときに、「若者と世帯主の定理」って言ってきた学生がいたんですね。何を訳したらこんな風になったのかなと思ったんですけども、これこそMDSのこの話で、ヤングーハウスホルダーの定理なんです。YとH、大文字だったんですけど、彼は多分辞書引いたんだと思います(笑)。

さて、ヤングさんとハウスホルダーさんが作った定理がありまして、これによってどういう条件であればMDSができる、あるいは空間行列の分析結果を距離とみなすことができるかという必要条件が明らかになっています。これについての資料としては、千野先生という愛知教育大学の先生のウェブサイト「多次元尺度法講義ノート」というのがあって、そこに丁寧に数式展開も含めていろんな説明がありますので、それ見ていただいたらいいんですけど(千野,2003)。大変なので簡単にかいっつまんで話しますと、要は座標と座標のベクトルが組み合わさったものが、どういう理屈で距離行列になっているか、という行列の計算についての話です。こういう風に要素間の座標を考えれば、距離とみなすことができるよというようなことを述べた定理なの

ですね。数学的な細かいところに入ってもいいのですが、大雑把に言いますと、基本的なメカニズムは正方行列、距離で表現されているサイズ N の正方行列、 $N \times N$ の行列を考えます。それを A と A' (A を転置させたもの)、同じ行列を二つ掛け合わせたものの形に分解することができれば、 A という行列の中の要素がそれぞれ座標を表していると解釈できるよ、っていう仕組みです。この時、 A の行列のサイズは m 列しかなくて、数式では $m \ll N$ って書きます。要するに m は N より随分小さい、十分に小さいという意味です。多次元尺度構成法で地図を描くという場合は、大体、二次元とか三次元なので、 m は 2 とか 3 とかっていう数字が入るのですが、二次元とか三次元ぐらいのサイズで分解できるとすると、 A の方の行列の要素は座標になっていると見なすことができますよ、というようなことなんだと思ってください。

正方行列をサイズの小さな矩形行列の積の形に分解できれば座標とみなせるよという話で、そんな都合のいい行列の分解なんかできるのかと、パッと見、そう思ってしまうかもしれません。でもこれ、数学的なトリックとはいいませんけれども、数学的な技術を使うとできるんです。固有値分解というんですが、この辺は線形代数の話になりますので、手前味噌ですが拙著「言葉と数式で理解する多変量解析入門」(小杉,2019)を読んでいただければ、イメージをつかんでいただけるかとおもいます。ひとまず皆さんにはこの行列の分解という話のエッセンスだけつかんでいただければというふうに思います(Figure 7)。



- ・サイズ N の正方行列を $M \ll N$ な 2 つの $N \times M$ 行列の積の形に分解できると、その行列の要素は座標になっているとみなせる

Figure 7 行列の分解についてのスライドより

さて、データを距離行列にすることなんですけれども、普通のデータは矩形の、縦長いデータが多いんじゃないかなと思います。列に変数が入っていて行にケースが入っている、一人の人に問い 1、問い 2、問い 3、このあたりの次元で評価してもらったというようなのが一般的な調査系のデータだったりします。でもこれ、もちろん縦は一行だけでもいいんです。そこから距離行列を構成しますから。最初からペアについての、対刺激のデータでもいいんですが、一般的なローデータについて、列変数同士の 1 列目と 2 列目の差の二乗を足したやつをルートしたのを作ると、変数 1 と変数 2 の類似度行列というのができます。それと同じように 1 と 2、1 と 3、1 と 4、2 と 3、2 と 4 という

風にやっていると、変数の方を使った下三角行列の距離の行列というのを作ることができます。これは別に列方向・変数の類似度に限ったことではなくて、1 行目、2 行目、3 行目、つまり一人目の人のいろんな反応と二人目の人のいろんな反応の、同じ変数同士の差の二乗を取って足したやつをルートすれば、ケース同士の距離を考えた下三角行列の計算ができるわけです。

これ、どちらの側からアプローチしても構いません。変数を分類だとか、マッピングしたいっていうのであれば、変数の側の、だいたいサイズ M ぐらい、ちょっと小さめの列になるかと。まあせいぜい 20, 30 ぐらいかなと思うのですが、それぐらいの 20×20 とか、 30×30 ぐらいの行列ができますし、100 人、200 人分のデータがあって、一人一人の違いが見たいと言うのであれば、 100×100 、 200×200 みたいな距離行列、いずれにしても最近の統計ソフト R でしたら `dist` 関数を使っていただくと一瞬で距離行列が作れますから、そんな感じでどちら側についても距離行列っていうのを考えることができるわけです。もちろん相関行列みたいなのを距離とみなすとしていただいても構いませんが、とにかくこういう対称の行列ができます。対称の行列ができれば先ほどやった行列の分解の数学的な技術を使いますと、少ない次元に対象なりケースなりを布置した座標というのが計算されることになるわけです(Figure 8)。

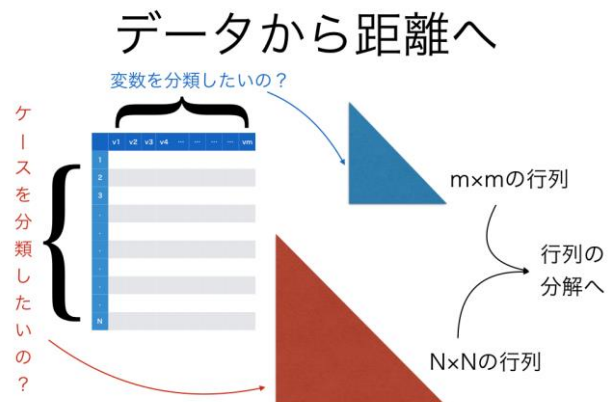


Figure 8. データから距離を作る。スライドより

実は多変量解析の世界はいろんな計算方法、分析方法がありますが、やっていることはすべて先ほどお示ししました、この行列をうまく分解してより少ない次元の形に分解して表現すると、同じ A と A を転置したやつの積和の形で作られる行列に分解するということをやっているに過ぎません。そのもとになるデータだとか、分解の仕方が違うときに分析名が変わっているわけです。例えば元々の関係、関連を表す行列、これは先ほども言いましたが、行の側でも列の側でも何でも構わないんですが、類似度なり非類似度なりを表している、関係を表す行列があったとします。たとえばそれが、変数同士の相関係数であれば、それを固有値分解、行列を分解する分析方法が因子分析と呼ば

れる手法なわけですが。相関行列ではなくて、平均とかの情報を含んだままの分散と共分散からなる分散共分散行列であれば、これを固有値分解をして得られる答えは主成分分析(PCA ; principle component analysis)だったり、構造方程式モデリングの中の潜在変数を決めたりするあの SEM になります。SEM は複雑なモデルなので、そのベクトルの中に数式がいろいろ入ってくるんですけども、分散共分散行列を元に少数次元にまとめようというところは同じなんです。元の行列が距離行列であれば、今日お話ししている MDS をやったことにもなります。あるいはクラスター分析というやつも、基本は距離行列からスタートして分類するんです。というか、同じことです。見方が違うだけなんです。他にも関係、関連を表す行列というのはありまして、例えばクロス集計表は、設問に「はい」か「いいえ」で答えたのを、男女で分けて集計した 2 かける 2 のクロス集計表とかがあっていうのがあったりしますが、ああいうやつが分析対象です。テキストマイニングと言われる分析方法では、内部でそれをやっています。テキストマイニングというのは、形態素解析と多変量解析の合わせ技の総称なんです。基本的には自然言語で書かれたテキストを、まず形態素解析というもので言葉の単位に分解しまして、品詞あるいは単語ごとに単語 A と単語 B が同時に使われてる回数は何回かというような行列を作るわけです(共頻度行列といったりします)。サイズは非常に大きくなりますけど、正方行列です。これを同じように分解することで、どういう言葉がたくさん使われますよ、というマッピングができるわけです。よく使われる多変量解析の手法としては双対尺度法(数量化 III 類)というのがありますが、その共頻関係をそのまま距離とみなして MDS をかけたりクラスター分析をかけたりしても構わないんです。テキストマイニングのツールで、デフォルトで MDS とかクラスター分析が入ってる KH Coder というフリーソフトウェアをご存知の方もいらっしゃるかもしれません(Figure 9)。

要は、ここで勘所としてつかんでいただきたいのは、MDS っていうのは特殊な分析方法ではなくて、他の分析方法と同じようなことやってるんだということです。因子分析法とか構造方程式モデリングのように各変数、データがどういう風に作られてきたかという細かいメカニズムを考えるのではなくて、座標でプロットするだけですからデータに対しても自由度が多いですし、そこから出てくる結果の読み取り方も様々に楽しむことができます。これが MDS のメリットなのだと思っておいください。

データの相と元 MDS の話をするときはそのデータの性質について議論されることがありますので、ちょっとこれをうちく程度にお話ししておきたいと思います。データというのは世の中にはいろいろありますが、区別の仕方として相と元という考え方があります。英語では mode と way になります。

距離と分析法

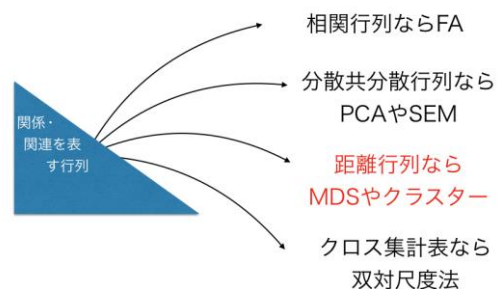


Figure 9. 距離と分析法に関するスライドより

最頻値も mode と言いますが、それとは違ってデータの mode という話をします。データの mode (相) とは一体何かというと、何種類の変数の組み合わせからできていますか、ということを表すものです。way (元) というのはそれを何回組み合わせましたかという意味です。一般的に一人ひとりの個々人が、ある変数のセットに答えましたよというデータは、変数と個人という 2 種類の情報が組み合わさって出てきています。組み合わせた回数が個人で一回、変数で一回なのでこういうのを 2×2 の 2mode-2way のデータと言ったりします。同じ人に対して、同じ変数を使って春・夏・秋・冬のような感じで、複数回調査しましたよいうことであれば、含まれるのは変数と個人、これを 3 回繰り返しますから 2-mode-3way のデータと言ったりします。なんでこんなことを区別するのかと言うと、例えば車の販売台数みたいなことからトヨタとホンダ、どちらの販売台数が多いか、何台分ぐらい離れてるかとかっていうデータを考えるのであれば、種類が車の売り上げだけになったりするので、1mode-2way のデータになるんですね。MDS の場合はこういうデータになることがあるんです。私が見た中で一番組み合わせが大きかったのは、社会調査のデータなんですけど、誰が何チャンネルの何時台のどういうジャンルのテレビ番組を見たかっていう 4 相データですね。どういう種類のデータを取ったかということによって、このデータの特徴が変わってくるんです。1 相、2 相ぐらいのデータであれば、いわゆる多変量解析みたいな大体使えるのですが、3 相、4 相ぐらいの多変量解析になると、それ用の分析モデルを適用することになります。このあと 2 相 3 元の MDS をご紹介しますが、mode と way の数によっては分析方法もちょっと変わってくることもある、ということは頭の片隅に置いていただければと思います。mode がたくさんになってきても、最近では多変量解析的な技術で潜在変数として扱うこともありますし、これをわざと潰してしまうこともあります。2 相 3 元のデータを 2 相 2 元に落として類似度行列を作る、というようなこともあります。

計量/非計量 MDS の違い

さて前半の最後に、MDS の種類についてもうひとつ大きな違いがありますので、それをちょっとお話ししておきたいと思ひます。

まず用語の説明から。MDS はいろんな対象を使って地図を描きます。この対象を使って作られた座標で、座標によって作られた地図のことを特に配置(configuration)と言ひますので、この用語を知っておいてください。もうひとつ、計量 MDS(メトリック MDS)というのと、非計量 MDS(ノンメトリック MDS)という、二種類の MDS があることを知っておいてください。メトリック云々って何だという話ですが、これはデータの性質に由来する区別です。

古典的な MDS は、ヤングとハウスホルダーの定理、MDS がどうやってできるのかという必要十分条件を示したあの定理のおかげで、地盤が固まったんです。あの定理そのものは、実はある対象を中心に行列を分解するという、そういう条件でやれば解けるっていう話なのですが、一般に MDS では誤差が入ってくるようなデータですから、特定の対象を中心にしてそこからの距離みたいな考え方は不自然だということで、行方向にも列方向にも平均値を取って標準化して原点を中心にした地図を描くという方法をとります。ですので原点を中心にした配置が得られるんですね。最初のヨーロッパの例でお示しましたように、東西南北の向きはぐるぐる変わります。変わるというか不定なので、自分で見たいように見ていただいたらいいんですが、いわゆる因子分析のように軸の解釈をして、右方向に行けば行くほどどうで、左方向に行けば行くほどどうでという解釈はできないのが普通です。どうしても置きようがありますから、反転、回転の可能性があるというのは知っておいてください。もちろん、デフォルトで出されたものを、なんらかの説明率が最も高くなるように、例えば分散が最も多くなるように座標軸を回転するバリマックス回転ということも可能ですし、何かターゲット行列を作ってそれに合うようにプロクラステス回転して軸を止めるという風なことも可能ですが、そういった外部的な基準がない場合は、原点に対して、回転に対して不定ですので軸の解釈なんかはしないように注意していただければと思います。

話を戻します。この古典的な MDS はメトリックな MDS といいます。ではノンメトリックというのは何かというと、もともとのメトリックな MDS というやつが行列の数値的な分解をもとに座標を作りますよという話であつたように、ヤングとハウスホルダーがサポートした計量 MDS は基本的にはデータが比率尺度水準で得られる必要があるんです。先ほどの例で eurodist, ヨーロッパの距離データというのを元にしましたが、測定上の誤差はあるかもしれませんが、基本的には距離データですから比率尺度でいいだろうということで先ほどそれを例にしました。それに対して、人の判断、類似度の判断とか、類似度評定みたいなのは本当に量的な、メトリックなスケールになつて

いるのかと言われると、非常に怪しいですね。人文社会科学系のデータの場合はそこまで厳密に測定できているわけではないんじゃないでしょうか。そんなときに、対象 j と k の類似度を例えば δ_{jk} とおいて、それを何か次元の方にプロットするとしましょう。プロットしたときの距離 $d(j,k)$ というのを考えたときに、この類似度が例えば $d(j,k)$ と $d(l,m)$ を比べたら $d(j,k)$ の方が大きいとします。そこでプロットされる空間の方でも、この $d(j,k)$ の方が $d(l,m)$ よりも距離が近いというように、この大小関係だけを保存するというルールにして、それでプロットされる座標を決めていくとします。このように、厳密に比率尺度的な条件を抜きにして、類似度が近ければ対象のプロットされる配置の距離も近いという順序関係だけを使ってプロットできないかというように考えられてできあがったのが、非計量 MDS というやつです(Figure 10)。ですので被験者の判断が、例えば「綾鷹」と「いろはす」は決して似ているといえないけれども、「綾鷹」と「おいお茶」は非常に似ている、というような判断をだつたとしますよね。「決して似ているとはいえない」に 1 点、「非常に似ている」に 7 点とかつけたとしても、別に 7 点が 1 点の 7 倍も似ているとまでは厳密に言えない、というのは心理学を学ばなくてもわかるころだと思います。こういうときは、もう順序性しかないと仮定して分析できる MDS、つまり非計量 MDS を使うことになります。これくらいデータに対する要請が少ない方法なので、心理学的なデータなんかの場合は非計量 MDS の方が使いやすいと思います。他の多変量解析に比べて、ゆるい尺度水準に早くから対応していたというのが MDS の強みでもあるかと思っています。

非計量的 MDS

- ・ 計量 MDS は距離データが比率尺度水準で得られている必要があるが、人文社会科学系のデータの場合はそこまで厳密に測定できている場合は少ない
- ・ そこで対象 j と k の類似度を δ_{jk} , 多次元空間での距離を d_{jk} としたとき、 $\delta_{jk} > \delta_{lm}$ ならば $d_{jk} \leq d_{lm}$
- ・ となるように配置をきめる、というルールに弱める。
- ・ 推定には不適合度 Stress を目安にする (Kruskal の方法)

$$Stress = \sqrt{\frac{\sum \sum_{j \neq k} (\hat{d}_{jk} - d_{jk})^2}{\sum \sum d_{jk}^2}}$$

Figure 10. 非計量 MDS についてのスライドより

非計量的な MDS をやって、じゃあどうやってその $d(j,k)$ みたいな座標を定めていくんだという話になりますが、これは計算上の最適化アルゴリズムというのがいろいろ考えられていて、クルスカールさんという人が対象間の本当の位置と計量された位置の、差の二乗の平均とったやつを指標にしようと言ってます。これをストレス値といいいますが、これが一番少な

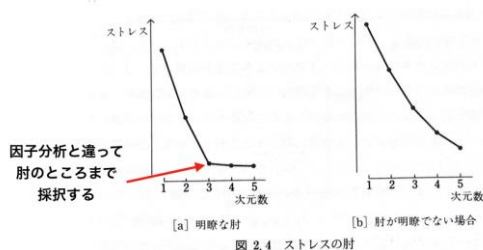
くなるように座標を探していきましょうというアルゴリズムが考えられています。このストレス値というやつですけれども、これは提唱者のクルスカルさんの基準で、小さければ小さいほどその再現性がびったりだということで、ストレスが 0.025 ぐらいだったらエクセレント, 0.05 だったらグッド, 0.1 だったらまずまずで, 0.2 だったら残念, という、一応の目安があったりしますので、これを見て当てはまりの程度はどうだったかを考えていただくことになります。

ちなみに因子分析とかと同じで、表現する次元が多ければ多いほど当然のことながら当てはまりは良くなっていくわけです。当てはまりは良くなっていくのですけれども、ここから先は誤差だよなとか、ここから先はそんなにストレスが大きく改善されないよなっていうようなところを探しながら、この地図のは、この類似度のデータは何次元ぐらいで表現したら良いだろうかという風なことを探していただくことになります。

ストレスの目安ですが、因子分析において固有値の減衰状況を見て因子数を定めるとかっていうやり方があるように、ストレス値を一次元のとき、二次元のとき、三次元のときと各次元で計算して、これの減衰状況を見て、ここまでとかっていう風にして決めましょうということになりますので、少しご紹介しておきます(Figure 11)。

Stressの目安

- Stressをスクリープロットのように並べて判断することもある



岡太・今泉「パソコン多次元尺度構成法」共立出版 より

Figure 11. ストレス値の目安についてのスライドより

右側にあるのがストレスの減衰がなだらかで、何次元にしてもいいかははっきりとした区分がないという状況です。因子分析なんかの場合でも、固有値の減衰状況で右側のような状態になっていると何因子構造かというのがはっきりつかめないとかっていう話になるんですけれども、MDS のときでも左側のように、ここでもうこれ以下はそんなにストレスは変わらないよとっていうのが出てきたら、そこまでの次元を採択しましょうというルールになっています。因子分析のいわゆる固有値の減衰状況と、切るところが違いますのでちょっと注意してください。因子分析の場合は、例えばスクリープロットが左側のようになっていたら二因子構造を採択するっていうのが基本的な考え

方ですね。この三因子目以降はもうほとんど誤差のようなものだから二因子で説明が終わりって考えるんですけど、ストレスの減衰で考える場合はその下がりきったところまで、左側の例でしたら 3 次元目まで採択するという考え方をしますので、そこだけご留意いただければなというふうに思います。

ちなみにここにも書きましたが、岡太先生と今泉先生が共立出版で出されてる『パソコン多次元尺度構成法』(岡太・今泉,1994)という本がありまして。これももうひょっとしたら絶版なのかもしれないですけど。これもとても良い本ですのでご紹介しておきます。この本はクルスカルの方法だとか、今日ご紹介するいろんな応用的な手法も入っていますし、因子分析法とかクラスター分析とか、コンジョイント分析なんかも入っています。一見バラバラな分析法が、なぜこの一冊の本に入っているのかと言うと、ちょっと前にもお話ししましたように、これらいずれの分析方法も似ているか・似ていないかという類似度を元に分類なり布置を求めるというような作業をしているという意味で同じだからです。そういう意味で一冊にまとめられるのですね。このパソコンっていうのを堂々と銘打っているところからお察しいただけるかと思うんですが、多分当時はこの「パソコン」っていうのがすごくて、前の方を見ると PC 9801 で実行できるソフトウェアがありますとって紹介があるんです。残念ながら、今のご時世だと何のことかわからないですね。できれば R か何かでパッケージ出していほしいところですが。中に書いてあることは今でもなるほどな、と思わされる、いい説明がいろいろ載っていますので、是非ご参考にしていただきたいと思います。

非計量 MDS の例

ではちょっと伴走サイトの方を見ていただけますでしょうか。具体的に非計量 MDS の例を示してみたいと思います。サイトの方に、M1 グランプリの採点結果より、というセクションがあるかと。みなさん、M1 グランプリはお好きでしょうか。僕大好きなんですけど。去年度も 10 人の漫才師が漫才をしまして、7 人の審査員が 100 点満点で採点しています。漫才がうまかった人が優勝するレースなんです。その採点基準についてですね、なんか言われたりするんですね。好みの問題だと発言したら、好みで判断されたら困るとかっていうような意見があったりするんですけども。でも、おそらくこの審査員たちはそれぞれ舞台上立つ人ですから、その人たちなりの感覚で、技術点だとか、テンポだとか、動きだとか、何かしらその人が大事にしているいろんな評価次元があつて、それらを踏まえて 100 点満点で言うところだ、という形にして表に出すわけですね。それがどういう次元で判断されていたのか、あるいは判断された人達はその辺が似ていて、どの辺が似ていなかったのかっていうことを MDS はやってくれるわけですから、今回のセミナーにおいても格好の材料かとも思います。このコードの部分よろしけ

ればRの画面にコピーしてちょっと実行してみてください。

伴走サイトの方にM-1 2018というデータセットがあって、点数のデータからプレイヤーの、漫才師たちの類似度行列ができます。非計量的なMDSをするときは、MASSパッケージの中に入っているisoMDSという関数を使います。このisoMDSという関数を使いますと、非計量的な、つまり順序性だけを担保した分析ができます。データから距離行列を計算していますが、審査員同士の距離が近かった、低かったという、類似度の順序関係だけを使って分析しますのでMDSなんかは非計量的の方が向いてるかと思います。このisoMDSという関数で、距離行列を与えてそれを何次元で表現するかを関数の引数で与えてやると、機械がさっと計算をして最適な布置を探して最終的にストレス値がこれくらいでしたよという答えを出します。えーと2次元を指定してやると、finalvalueで0.013671って出力されます。0.01だったら先ほどのクルスカルの基準でいっても0.01っていうのは素晴らしいと完璧の間になるのですが、これ実はなぜかこのisoMDSは百分率の形で出てます。ですからこれ0.01じゃなくて、本当はさらにその1/100の0.00013671のことなので、もうとにかくほぼ完璧にすごいぐらい二次元でプロットできるということがわかります。これ今は皆さんやらなくてもいいですが、次元数を探索するときはisoMDSの関数に次元数をfor関数か何かでぐるぐる追加してやって、何次元ぐらいがいいかな、というのを試してそれをプロットするというような形で見てくださいね。これ見ていただくとストレス値がもう二次元のところでもほぼストーンと落ち切っていますので、先ほども言いましたように、ここまで採択して二次元であるプレイヤーたちをプロットすればいいんじゃないかということがわかります。サイトの方にggplotパッケージを使ってプロットするコードもありますので、よろしければぜひお試しください。それをやるとこんなふうに出てくるかと思います(Figure 12)。

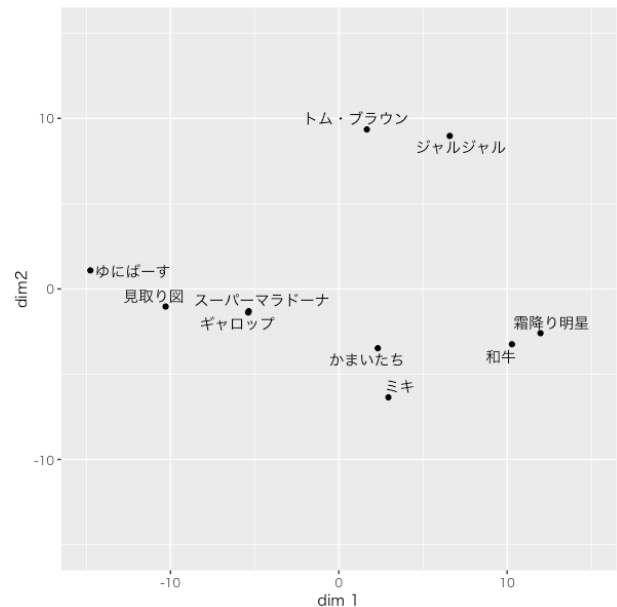


Figure 12. 伴走サイトより、M1 データの分析例

右上の方にトム・ブラウンとジャルジャルがいます。右下の方にミキと和牛と霜降り明星とかいて左の方にユニバースとか見取り図とかかいてですね、皆さん、「どんなネタだったかな」とか「どんな芸人だったかな」って思い出していただくと、何となくわかりいただけるかと思います。ちなみに僕だったらたとえばこんなふうに考えるんじゃないかなと、少し軸を回転させてみました(Figure 13)。

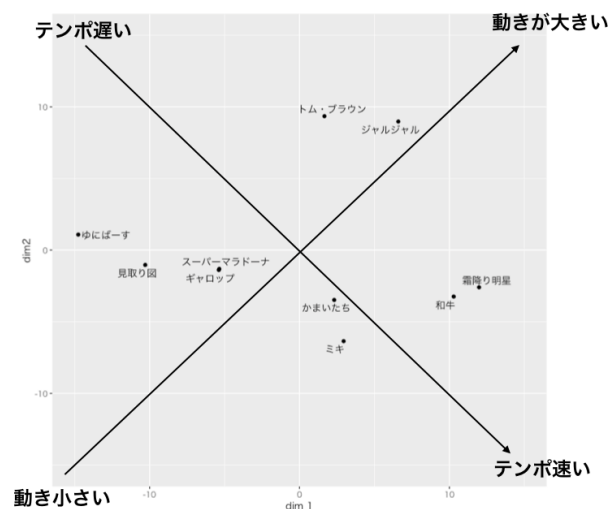


Figure 13. 軸を回転して考える

原点を中心にぐるぐる回りますので、原点を中心に大まかに回転させた方がいいかなと思って軸を描いたのですが、このように見ると、右下のグループはテンポが早くしゃべるミキとか和牛がありますね。ユニバースはすこしゆっくり喋りますから、左上はテンポはゆっくりめかなとか。トム・ブラウンとかジャル

ジャルは、動きが非常に大きかったりしますから、右上が大きく動く方で左下側はそんなに動かない方かな、といった解釈ができますね。いやいや、ジャルジャルの一本目はそんなに動きがなかったぞ、という反論もあるかもしれません。その辺の解釈は、いろいろあり得ます。もっとも、繰り返しますが、軸は回転しますので、最初に出されたこのときの一次元目と二次元目というのをそのまま解釈にあてるのは危険です。まずは自分で布置を見ながら、大体どの領域にどんなのがあるのかなっていうように見てください。

解釈は本当に自由にできますけれども、例えば類似度データ以外に対象に布置された情報があつたとします。例えば審査員の点数以外に街の声としてこの評価が高いとか、各コンビニの顧客動員数がどうだといった、他の説明変数があれば、座標の値を従属変数にして回帰分析を行い、どういう説明変数がこの軸を説明するか検証することもできます。解釈次元の妥当性を考える1つのやり方ですね。この方法はまたまた手前味噌ですが、「緑の紐本」の愛称がある小杉ら(2018)の中に、多次元尺度構成法の章がありますので、ご参照ください(Stalans, L. J., 1994)。

もともとが距離のデータなので、これにクラスター分析を重ねて、ここが一つのクラスター、ここは二つ目のクラスターというように、グルーピングしていくこともできます。また今回は芸人さんの方をプロットしましたが、評定の行じゃなくて列の方の類似度行列を計算して分析することもできます(Figure 14)。そうすると、立川志らくと上沼恵美子の評定はかなり離れているんだな、ということが見えたりします。点数をつけるときにどんなことを考えたかわかりませんが、離れているんだということは確かなわけで。中川礼二はこの二人の中間ぐらいだな、変なところにいるな、とかですね。富澤たけし、オール巨人、塙、松本人志の辺りがすごく一箇所にまとまるので、これはこれで彼らは何か共通の次元で判断してたんではないかなというようなことを考えたりすることもできます。たった一回の番組の7かける10ぐらいのデータでもこんな風にして遊べますので、ぜひいろいろ試していただければというふうに思います。

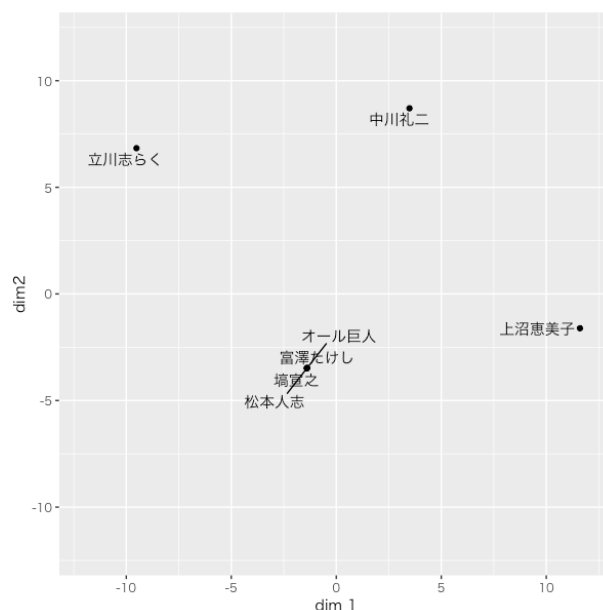


Figure 14. 審査員の方のMDS結果

前半を終えて

どどーっと喋ってきましたのでちょっと休憩を挟んでもいいでしょうか。何かこままでのところでご質問とか、ご意見ありますでしょうか。

質問：最終的に研究者の主観によって結果を判断するのか そうです。そもそもそういう意味で探索的な分析方法だと思ってください。潜在的な評価尺度次元を探すという技なので、本当かどうかは分からない世界です。

質問：次元数の探索によって四次元以上になることはあるのか 今までの経験では、ほとんどないです。一番多くても3とか4ぐらいで。もちろん変数の数とかにもよるんですが、2, 3次元ぐらいで収まることの方が多いです。僕の経験上ですが。

質問：4次元の場合はどのように分析・解釈するのか そのときはもう1次元目と2次元目のプロット、1次元目と3次元目のプロット、2次元目と3次元目のプロットみたいに、いくつもプロットを描いて、こっちの次元で見るとここが近いな、と試行錯誤しながら総合的に判断することになります。イメージするだけで難しいんですけど、実際経験上そんなにたくさんの次元が、十何次元出てきて困っちゃったってことは僕はあんまりないです。

質問：元データと相関データのどちらを使って分析したほうがいいのか どちらでも良いです。相関行列を距離とみなす場合は1 マイナス相関行列の絶対値、みたくに変形する必要があります。元データとどっちがいいのかっていうことですけど、これがMDSの弱点のところなんですけど、どちらでもいいんです。いいんですけど、結果はそれぞれ違ってきちゃうんですね。相関データの方が個人行なり列なりで標準化されていて、相対的な比較になっているので、単位が整っているという長所がありま

す。それに対して、元データはやっぱり単位こそバラバラかもしれないけど、測定値としてははっきりしているので、そういう意味のある違いなんだ、ということが出来ます。つまりどちらの良さもあって、どちらの良さと距離行列を作っていたでもいいです。ただ、それぞれの方法でやった距離行列から出てくる地図って全部違うんですよ。どれが正解かって言われると難しいところなので、解釈可能性みたいなのを考えていろいろ応用していただければと思います。

質問：ストレス値がなだらかな場合はどうするか 解釈しやすい次元に定める、というのが1つの基準です。数量化Ⅲ類のときなんかもそうですけど、次元が増えてたくさんの次元になるとやっぱり分かりにくくなるので、最初からもう2次元で決めちゃう、ストレスが良かろうが悪かろうが2次元で決めちゃうみたいな考え方も一つにはあります。ただその時は、ストレス値が良くないから暫定的な結果だ、大体の傾向だ、というようなことも踏まえて判断にいただいた方がいいと思います。

MDS の応用

では後半のお話に進みたいと思います。後半は応用的な MDS のお話です。三つほど話題を持って参りました。ひとつが個人差を表現する INDSCAL という方法、二つ目が地図に情報を書き加える方法で、三つ目が非対称関係に拡張する方法というやつです。順にご説明していきたいと思います。伴走サイト(小杉,2019)の方は上のタブのところでは応用とあるところを選んでいただくと、それらのコードが出てきますので見てください。

個人差 MDS (INDSCAL)

さて個人差 MDS からご紹介します。これは INDSCAL という名前がついています。D は発音しませんのでインスカル、というふうに読みます。individual な特徴を残したスケーリングということで ind- というのがついています。これは2相3元のデータ、つまり類似度評定のデータが三次元にあるようなデータを使うやつで、個人的にはすごく好きなモデルです。

どういうモデルかと言いますと、個人 i が対象 j と k の類似度を評定したということを考えたします。それが何らかの次元にプロットされます。そのプロットされるとき距離の計算は、いわゆるユークリッド距離の計算と同じですが、個々人の特徴はこの各ユークリッド距離に個人 i が t 次元目に持つウェイトをかけたものとして表現されるというのが個人差 MDS の考え方です(Figure 15)。

個人差MDS

- INDSCALという個人差を表現するモデルがあります

- 個人 i が対象 j と k の類似度を δ_{ijk} とつけたとします。

- 全ての個人の布置の原型となる共通布置空間では、

$$d_{jk} = \sqrt{\sum_{t=1}^p (x_{jt} - x_{kt})^2}$$

- 個人の特徴は次元に対する重みとして考えられて、

$$d_{jk} = \sqrt{\sum_{t=1}^p w_{it} (x_{jt} - x_{kt})^2}$$

Figure 15. INDSCAL についてのスライドより

数式だけで見ていてもピンとこないかもしれませんので、こんな例を持って参りました。これを見ていただきますと、プラスと三角と丸と四角がプロットされています。これらが対象ですね。その類似度評定の結果、りんごとマンゴーは似ていますか、マンゴーとみかんは似ていますか、みかんとメロンは似ていますか、といったデータを取って、赤の四角で描いた世界にプロットされたものだと思ってください(Figure 16)。

INDSCALのイメージ

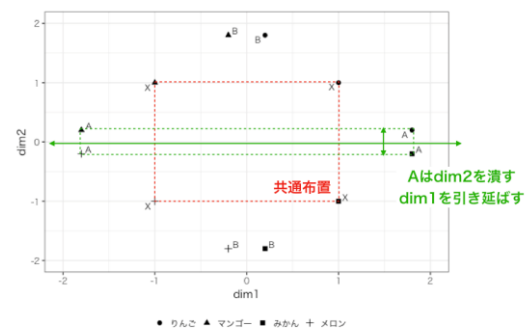


Figure 16. INDSCAL の分析結果のイメージを掴む
スライドより

そのりんご、マンゴー、みかん、メロン、それぞれの類似度調査を A さんと B さんの二人にしました。プロットされる軸が A さんと B さんとは歪んで知覚されている、軸の重要度が違うよという考え方、これが INDSCAL のイメージです。全体の共通空間としては赤字のように等間隔になっていたとしても、A さんの方の世界は横に伸びているわけです。これはなぜかということ、この第一軸の評価は、三角と四角ですからみかんとメロン、りんごとマンゴーというふたつですね。わからないですけど、歯ごたえの良さとみずみずしさみたいな次元があったのですかね、それを A さんは非常に重要なものだと思って評価し

ている。こういうときはこれが横にギュウッと伸びるわけです。これに対してマンゴーとメロン、リンゴとミカンですね、形の丸さを表す次元でしょうか、それについてはAさんはそんなに気にしないようです。このときは、この軸がキュッと縮んでいきます。Aさんは第二軸を縮めて第一軸を強調して表現する、そういう形で個人差を表現できるのです。同様に、二人目の人は第一軸はそんなに重視しないが、第二軸が大事だというようにプロットすることもできます。こうして軸の伸縮で個人差を表現する、というのがこのINDSCALの面白いところですね。個々人がどれくらい歪んでいるかを考えることもできますし、共通空間を考えて平均的にはこういう風になるんだと、いろんな評価の違いはあるけれども共通した世界はここなんだというプロットもできるという意味で、個人差の表現力が増した面白いモデルではないかなと思います。

具体的な計算のときは、重みがマイナスにならないように制約します。人によっては次元がひっくり返っちゃうと、ちょっと話がややこしいので、そうはならないようにします。あと、次元の重みが何倍してもよいということになりすぎたら困るので、重みのノルムは1にしておきましょう、といった制約を課すことになりますね。その制約のもとでプロットを描きますので、ここまでのお話にあった、地図の回転に対する自由度がなくなるのも特徴です。つまり共通次元はもう個人差みたいなのは軸の伸縮の率で表現されて、共通次元で説明ができるので、軸の解釈ができるというのがこの分析方法の特徴の一つでもあります。人によってはそこが歪んでるかもしれない、という可能性については、INDSCALの重み行列を対角じゃなくて、非対角要素にも大きさがあるもの、つまり軸の相関を許したような形で表現する、IDIOSCALという方法もありますので、こんな方法もあるんだって知っておいていただければと思います。

ちょっと具体例です。以前、豊田先生の『たのしいベイズモデリング』(豊田,2018)っていう本に、INDSCALのベイズ版を分担執筆で提供したのですが、これの例の一部を使って、このINDSCALの実例をお見せしたいと思います。どんなデータかと言いますと、観光都市が10個あげられていまして、札幌、飛騨高山、舞鶴、佐世保、志摩、秋吉台、野沢、道後、湯布院、宮古島という10個の町があって、この旅行先似ていますか、似ていませんかという調査をやったんですね。元々のデータは何百人ぶんかあるんですけども、今回はその一部、五人ぶんを取り出したデータを伴走サイトの方に掲載しています。これを使って試してみたいと思います。INDSCALはsmacofというパッケージでできますので、これを事前に準備しておいてください。ファイルを読み込んで、マトリックスを人数分足していきますから、Rでいうところのリスト型で10かける10の行列が5人分入っているリストのようなデータを作ります。INDSCALはsmacofパッケージにある、smacofIndDiffという関数を使うことができます。そのデータが間隔尺度水準なのか比率尺度水準なの

か、順序尺度水準なのかみたいなのも指定できますので、順序尺度水準でINDSCAL方法、つまり軸の対角線だけ重みづけるという風な形で分析しなさいという、あつという間に結果が出てきます(Figure 17)。

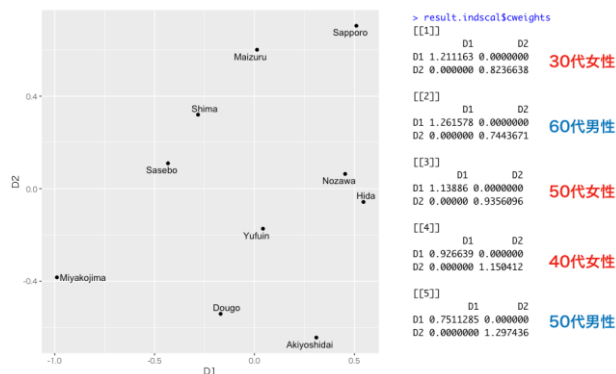


Figure 17. INDSCALの実践例

宮古島がこっちにあって、秋吉台がこっちに行って、札幌があっちに行ってという地図ですね。これ実際の地図じゃなくてイメージ、イメージですよ。観光地のイメージでが似ているかどうかというイメージのマップが出てきているわけです。これに加えてこの5人の人、一番目の人は30代女性ですが、この人は第一軸を1.2倍、ですからこれが横に伸びるわけです。つまり一番目のこの人にとっては佐世保と野沢の距離なんかはこの共通次元よりも遠くへ、ちょっと似てないという判断をされている。それに対して同じく一番目の人は縦軸が0.82倍されますから、上の方にある舞鶴と湯布院の距離がきゅっと、8割ぐらいに短くなって知覚されているよということがわかります。一番の30代女性はそうですが、一番下の50代男性なんかは逆に、横軸が0.75倍、縦軸が1.29倍ですから、この人は第二軸の方が似てない。いわば歪みを繊細に区別できるので、こうしたプロットができる。

こういう風に個人差を表せるようにモデルが展開しているのもあります。いろんな人に似てる似てないの判断をしてもらった後で、個人差をつぶすことなく分析したいとかっていうときはこのINDSCALなんかを使っていればよいんじゃないかなと思います。

地図に書き加える方法

二つ目の応用例です。MDSで描かれた地図というのは、対象間の類似度を反映した地図なので、その地図に例えば何かプラスアルファとして、追加で情報を書き込むとかっていうことだってしていいはずですし、できるよというお話です。地図に何らかの情報を追加することによって、表現力あげましようというような例がありますので、それを二つほどご紹介したいと思います。

一つは、この似ている・似ていないの判断に加えて、各対象が好きなどうかの判断を追加する、選好度の写像分析、Preference Mapping をご紹介します。同じように、類似度の世界の中の力関係みたいなのを表現しようという Abelson Map というのもありますので、この二つをちょっとご紹介したいと思います。

Preference Mapping 一つ目の Preference Mapping というやつですが、類似度評価をされた地図の世界の中に個々人の点を追記しようという考え方です。その点はいったい何なのかと言うと、その人の理想的な点だというわけですね。この理想点というのは何かと言うことですが、類似度評定に加えて、それぞれの対象の好き嫌いの情報を追加で収集します。そうすると、似ているか似ていないかというこういう地図の世界だけけど、私の一番好きな味はこの地図でいうとこの辺りにはあるはずという、そういう個々人の好みの理想的な場所を求めるための点が描けます。その理想的な点と、その地図で描かれているそれぞれの対象との距離が、選好の強さを反映します。原理的には、個人の理想点の座標を何か一つ考えたとして、その座標点が対象との距離の一次関数で表現できると考えるわけですね。一次関数で考えますから、要はこの選好度 s_{ij} 、i さんの j に対する好意の度合いというのは、距離の二乗かける個人の定数+誤差みたいな形で表現したらいいじゃないかというわけです。つまり、こんな形にすると回帰分析で解けるので比較的簡単にできます(Figure 18)。

Preference Mapping

- ・ 類似度空間の中に「個人の点」を追加する
- ・ この個人の点は**理想点**を表す=理想点と対象の点との距離が選好度を表す
- ・ 個人iの理想点座標を y_{it} 対象jに対する選好度を s_{ij} とし
 このように表す $s_{ij} = a_i d_{ij}^2 + b_i + e_{ij}$ ← 回帰分析で解ける
 ただし $d_{ij}^2 = \sum_{t=1}^p (y_{it} - x_{jt})^2$

Figure 18. Preference Mapping についてのスライドより

そのときの距離っていうのは何かって言うと、対象の布置で得られた点と個人の理想点との距離になっているわけです。この数式だけから展開していくと、最終的にこういう形になります(Figure 19)、上の方の式の右辺を見ていただきますと、座標値の二乗と座標と誤差の形に分かれているように分解できるのですね。MDS では対象の座標というのが計算で出ますから、この座標が出たら、座標にこの数式の展開をした回帰分析をして、その回帰係数から個人の理想点を算出するということができます。この辺の数式の展開とか導入に関しては、先ほどご紹介しま

したこの『パソコン多次元尺度構成法』(岡太・今泉,1994)の8章か9章の後半の方に書いてありますので見てみて下さい。

Preference Mapping

- ・ 先ほどの式を展開すると

$$s_{ij} = a_i \sum_{t=1}^p x_{jt}^2 + \sum_{t=1}^p \beta_{it} x_{jt} + c_i + e_{ij}$$

- ・ MDSで対象の座標(x_{jt})は元待っているから、回帰分析を行なって回帰係数を算出し、次の式で理想点を求められる。

$$y_{it} = -\frac{1}{2}(\beta_{it}/a_i)$$

詳しくは岡太・今泉「パソコン多次元尺度構成法」共立出版を

Figure 19. Preference Mapping の計算方法に関するスライドより

ちょっと具体例です。先ほど M-1 の審査でやった、距離の世界があります。それに自分なりの点数をつけます。私個人の感想です。皆さんご自身の好きな数字を書いていただいて結構です。私は個人的には霜降り明星はそんなに良くなかったと思いますし、私はもっとジャルジャルに行って欲しかったので、そういう評定値を書いてみます。その上で、残念ながらこの Prefmap は R のパッケージで提供されるような関数になっていませんので、伴走サイトの方にコードを私が書きました。と言っても、回帰分析をやってこの係数を読み取るだけなので、大したことはないんですが。まず先ほどの二次元布置を見て座標が出てきます。この座標を使って各次元を二乗して足した項が必要なので、それを計算します。それと一次元目の数字と二次元目の数字でもって、自身のスコア、つまり Pref って書いたところが、私が各プレイヤーにつけた選好のスコアなんですけど、その情報を追加しましてそれを回帰分析します。そうすると二乗の項にはこの数字、一次元目の係数はこれで、二次元目の係数はこれみたいなのが出てくるので、これを使って座標を-1/2 かけるベータ分のアルファみたいなので計算して出したのがこれです(Figure 20)。ここに私の理想点があるわけです。

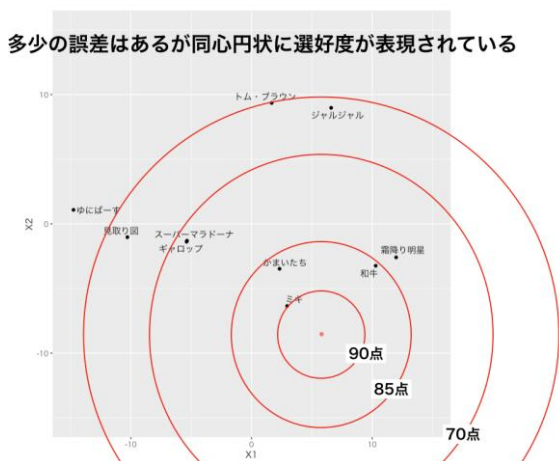


Figure 20. Prefmap の結果の例 (スライドより)

これが何なのか、という話ですが、これを見ていただくと、要はこういうことですね。あそこを中心に対象との距離が私の選好の度合いを表していますから、僕、多分ミキが一番好きなんです。それよりちょっと遠いところにかまいたちとか和牛とか霜降り明星とかがいて、だいたい僕のつけたスコアでは 85 点ぐらいのラインで、70 点ぐらいのラインが外にあってと、こういう風に同心円を中心にして自分の選好度が表現されているようになるわけです。いやいやジャルジャルの評定値はもっと高かったぞ、という意見もあるかと思いますが、確かにそういうズレてるところは出てきます。10 組のトータルの中で誤差が一番少ないのがあそこというだけで、一個一個のケースを見たらずれていることももちろんあるんですが、それはあの地図を作った審査員たちの評価と僕の好みとのズレでもあるので、その辺はちょっとご容赦いただいて、少なくともこうやって表現はできるということです。

今回は類似度評定と私の選好のデータが別でした。つまり類似度評定はプロの漫才師達がやっていますけれども、別に私のこのスコアを元に対象同士の距離を計算して、地図も自分の好みで変えて理想点も自分の好みで描いていいんですよ。ともかく地図さえあればそこに点を追加できるよということを知っていただければと思うわけです。ちなみに、これは、小杉はあそこって風に点を打ちましたけれども、同じ地図に複数の人の点を打つことだってできます。もしこれが先ほどの例の飲み物の類似度評定とかだったりすると、どうも多くの人の理想点はこの辺にあるけれども、そこに類似した商品がなさそうだということであれば、なんか開発のヒントみたいなものになったりするかもしれないわけですね。そういう風にして使える個人の好みの評定マップというのもあるよということですね。一つご紹介でした。

Abelson Mapping つぎに、もう一つの Abelson Map という技をご紹介します。これは正式名称があるわけではないんですが、エイベルソンさんが 1954-55 年の論文で提唱したやつで。

これは各対象に個々人の強い思い、重みが載っているというふうに考えます。その重さに応じて、その点から対象からじわじわと力が溢れ出ていって、その力がどうやって拮抗するかというような、場の理論、電磁波の理論みたいなのを応用した描き方です。これで地図上の各点の力がぶつかっているところ、ぶつかっていないところみたいなのを計算して、心理空間を描きましようとかっていう技なんです。ちょっと見てもらった方がわかりやすいかと思いますので、先ほどの M-1 の地図にですね、自分の好みを追加するという例でやってみます。私先ほど 70 点、80 点とかをつけましたがその平均点を引いたようなスコア使っています。何でこういう風になっているかというと、0 をまたいでプラスマイナスに描いた方がこの後出てくる地図が見やすいからです。なので、平均点を引いた形で中心化しました。

Abelson Map のパッケージとかありませんので、恥ずかしながら伴走サイトに私の自作関数を書いてありますから、それを読み込んでいただいてプロットしていただければ、こんなような地図が描けるのかと思います(Figure 21)。

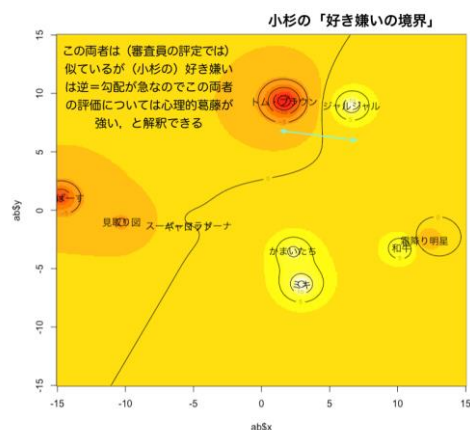


Figure 21. Abelson Mapping の例 (スライドより)

これ、白いところが山の部分です。赤い所が谷の部分です。ミキなんかは割と好きなので、ここから何かポジティブな力がじわじわと出ているわけです。ネガティブな、僕、多分トム・ブラウンが嫌いなので、あそこ谷があつてですね、間に線が走っているところ、これがプラスマイナスゼロのラインなわけですね。だから僕の好みは好き嫌いがこの境界で別れているということが読み取れるわけです。あの 0 のラインというのが私の好き嫌いのラインになっています。特にトム・ブラウンは底です。ジャルジャルは僕好きなので、山になっています。山と谷がすごく近くにあるというのはどういうことかという、類似度は似ているんですが好き嫌いが逆なので、個人的にはあんまり嬉しい状態ではないわけです。だから僕はこの状態にあると、なんかジャルジャルはもっと点数が高くないと、とか、トム・ブラウンはもうちょっと低くても良かったんじゃないかとか、なんでトム・ブラウンとジャルジャルを審査員と一緒にしちゃったんだ、と

いう心理的な葛藤が起こる場面という風に解釈することができるんじゃないかと思います。これは非常に、心理学的なフィールドをみたいなのを考えたモデル化をしています。電磁場のメタファーである、という仮定を受け入れるのであれば面白い絵の描き方ではないかかと思ひます。ちなみに、これはエイベルソンのやり方っていうので書いていますけれども、空間統計学のクリギングの手法を使って、実際のマーケティングデータにこの好き嫌いの勾配を書いているという例が朝野先生の本(朝野, 2010)なんかであつたりもしますので、ひょっとしたらまだまだ応用可能な、展開可能なジャンルなのかなとも思ひたりしています。

非対称 MDS

さて最後です。非対称多次元尺度構成法です。非対称っていうのはどういうことやねんという話ですね。そもそも距離は対称じゃなかったのということですが、距離の定義は対称なんですけれども、元々のデータが非対称なことっていうのは結構あるわけです。例えば好きな人に嫌われているという、片思いみたいな世界もある。ビールから発泡酒に変えたんです、プリン体のこと気にして変えたんですとかっていう人がいても、発泡酒からビールにはなかなか変える人がいない、とかっていうブランドスイッチングの問題ですとか。東京には人が集まっていますけれども東京から地方にはあまり人が流れていかないと、国際貿易で台北は黒字だけど台中は赤字だとか。そういうことを考えると、二者間で二つの対象間の関係が必ずしも対称になってないことっていうのは結構世の中にはあるわけです。この非対称情報に意味があるので、その情報を取り込みたいというときの MDS が非対称 MDS というやつです。心理学的なデータでは、先ほどのモース信号の間違いなんかも基本的に非対称ですね。A から B には間違いやすいけど B から A には間違いにくいとかってこともあつたりするので、そういう風な情報なんかも分析することができます。

さて、どうやってその非対称情報を分析するのかといったら、いろんなモデルがあるんです。一つはなかなか破天荒な方法でして、なにかというと、距離空間が非対称性を許さないんだつたら空間の方を拡張すればいいじゃない、という千野先生のやり方です。もう一個は、その距離は距離で普通にユークリッド空間に描いておいて、それに非対称情報を書き加えるというやり方でやったらいいじゃないかというやつです。ちなみに先ほどの Abelson Map も高さのところに非対称情報を描いて、非対称情報の山と谷みたいなので描くようなこともできます。それはその第 3 次元に外挿する方法なので、意味は何とでもつけようがあるというやり方なんです。ともかく、いろんなモデルが考えられていますので、2, 3 ご紹介します。

千野先生の HFM 一つ目の千野先生の HFM という技はエルミート空間モデル、エルミート形式モデルというやつでして、こ

れは何かって言うと空間の特徴の方を拡張しましょう。実空間でやっているからダメなので、複素空間でやってしまひましょうという技ですね。どうやるかと言うと、行列を対称な部分と対称じゃない部分に分割して、非対称な部分は虚数なんだ、として行列を組み上げると、その行列の、先ほど分解の式がありましたが、あの数学的なやつがですね、少なくとも答えが全部実数で得られる。そして固有ベクトルの方が、座標ですが、その内積が距離として定義できる空間になっているという技なんですね (Figure 22,23)。この式見るとなんか難しそうな気がしますが、実際の計算はたいしたことなくて。

Hermitian Form Model

- ・ 距離空間を複素空間に拡張すれば内積で距離を表せる
- ・ 行列を対称部と歪対称部に分解し、歪対称部を虚数にしたエルミート形式の距離行列を作ると、固有値が実数で得られることが証明できるし、固有ベクトルが布置を表す

$$S = s_{jk} = \frac{1}{2}(S + S') + \frac{1}{2}(S - S') = S_s + S_{sk}$$

$$H = S_s + iS_{sk}$$

歪対称部

Figure 22. HFM のスライドより

Hermitian Form Model

<pre>> # 非対称行列 > Asym [,1] [,2] [,3] [,4] [1,] 0 1 1 7 [2,] 1 0 1 7 [3,] 7 7 0 1 [4,] 1 1 7 0</pre>	=	<pre>> # 対称部 > (Asym+t(Asym))/2 [,1] [,2] [,3] [,4] [1,] 0 1 4 4 [2,] 1 0 4 4 [3,] 4 4 0 4 [4,] 4 4 4 0</pre>	+	<pre>> # 歪対称部 > (Asym-t(Asym))/2 [,1] [,2] [,3] [,4] [1,] 0 0 -3 3 [2,] 0 0 -3 3 [3,] 3 3 0 -3 [4,] -3 -3 3 0</pre>
---	---	---	---	---

↓

\$Hermitian

<pre>[,1] [,2] [,3] [,4] [1,] 0+0i 1+0i 4-3i 4+3i [2,] 1+0i 0+0i 4-3i 4+3i [3,] 4+3i 4+3i 0+0i 4-3i [4,] 4-3i 4-3i 4+3i 0+0i</pre>
--

Figure 23. エルミート形式への変換方法 (スライドより)

非対称な行列があつたとします。これ(Figure 23)見ていただくと、[1,3]の要素が 1 なんですけど[3,1]の要素が 7 だったりしてこれ非対称になってます。1 の人は 3 番の人は嫌いなんだけど 3 番の人は 1 番の人のことが好きとかそういうデータだと思ひていただいても結構です。これを転置して足して二で割つたやつ、転置して引いて 2 で割つたやつに分けます。それは何かと言うと、足して二で割つたら平均になりますから対称行列ができるわけです。引き算をするとその対称からずれているものが

出てくるわけです。引き算した方のやつを虚数の世界に移しますので、一つの行列として対称部分+非対称部分の情報が行列として描くことができます。これを先ほどの行列のやつと計算と同じで分解しちゃおう、というのが HFM です。ちなみにこのエルミート形式モデルもパッケージがあるわけじゃありませんので、私の自作関数を伴奏サイトに載せておりますのでそちらをご参照ください。そうすると先ほどの四人のデータだったらこんな感じで出てきます(Figure 24)。

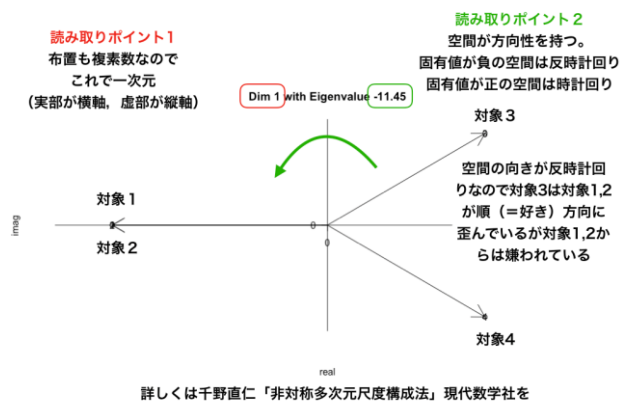


Figure 24. HFM の出力例 (スライドより)

これでどういう風に読むかという話ですが、ちょっと難しいのはですね、まずこれだけで一つの次元だということですね。2次元あるじゃないと思うかもしれませんが、これ実数と虚数になってるんですね。最終的にプロットされる座標が複素数で得られるので、これでも一次元です。複素数の1次元です。

もう一つは、出てくる固有値がマイナスになって-11.45 出ていますね。このマイナスになっている数値、プラスになったりすることもあるんですが、これはこの空間の方向を意味しています。固有値が正の場合は時計回りに、負の場合は反時計回りに向きが流れているというように読んでください。これはマイナスなので反時計回りに向きが向いているんですが、これはどう解釈するかと言うと、3は1のことが好きなんですけど逆方向は嫌い、と読みます。こういうようにしてさっきの四人のデータがプロットされているということになります。詳しくは千野先生ご自身のご著書でいろいろ書いてありますので、参照してください(千野,2003)。僕、個人的にはなかなか面白い技だなと思ってはいるんですが、複素数とかあんまり喜ばれる方が少なくでですね(笑)、結果も読みにくいというところもあるんではやらないのかもしれないです。

非対称情報を描き加えるモデル 情報書き足す系でいうと、岡太先生のモデルなんかはもう少し分かりやすいです。これは何かというと、そもそものプロットされる布置は MDS 空間ですから対称ですけども、これにこんな丸を書き足します。この丸の部分が非対称情報を表しています。j と k の距離、対称な部

分はあの長さですが、非対称な部分はどやって表すかという、具体的な距離はこの j が持っている楕円の端っこ部分から相手の中心を通して向こうの端まで、この長さが j から k への距離で、k から j の距離は、緑の字が間違ってますが、この k の外枠から出て j のこっちまで来ると。こういう風に表現することで、対称部分にまわりついている楕円の範囲がプラスとかマイナスとかで非対称関係を表現しようというモデルです(Figure 25)。こういう風にして情報書き加えるという技もあるわけですね。

岡太・今泉モデル

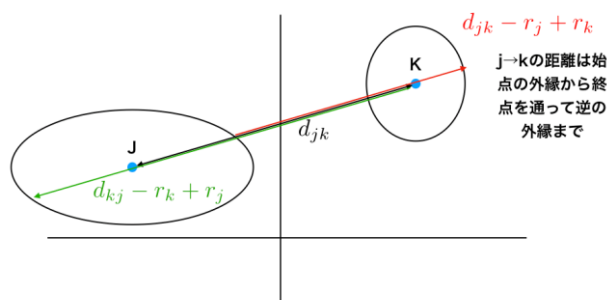
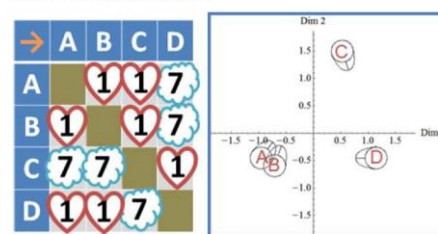


Figure 25. 岡太・今泉モデル (スライドより)

最後に、フォン・ミーゼス分布を使う、荘島先生の考えている面白い技があるんですが、これは下手な説明するよりもこちらのサイト見ていただいた方がわかりやすいかと思います。何かという対象に確率分布をまわして、だいたい全員がどっちの方向を向いているかというようなことを考えるモデルです。さっきの HFM と同じデータセットなんですが、これでやるとこの結果はこんなふうに出ます。A と B はなんとなく向こうの方が好きなんですけど、C はこっちの方を向いていて、D はこっちの方向向いてるみたいです、全員がどっちの方向に歪んでるかとかっていうのがプロットされる方法としてあります。

vMモデルの出力例

HFMの出力と比較してください



・この分析は荘島先生のサイトにあるアプリを使って実行することができます。興味がある人は是非！

Figure 26. 非対称フォン・ミーゼス分布モデルについて (スライドより)

なんでこの話したかと言いますと、非対称 MDS は実は世界の中でも日本が特に抜きん出ている技術でして、この研究者、岡太先生、今泉先生、千野先生、荘島先生というトップクラスがなぜか全員日本にいたので、日本では妙に非対称 MDS の研究が進んでいます。もし興味がある人は行動計量学会の中に非対称部会というのがありますので、非対称な人たちがやっている、研究が非常に楽しい世界ですので、ぜひご参加ください。

今日のまとめとしていいますと、覚えておいて欲しいのは、要は距離から地図を作る技術だということです。データがしっかりしていれば、それっぽい評価次元の地図っていうのを作ることができます。ですので、データの何を距離とみなすかというのが面白ければ皆さんの様々な領域で応用できるんじゃないかなと思います。しかも出てきたのはただの地図なので、それに落書きしたり色を塗ったり加筆修正したりとかってというのが結構自由にモデルを拡張できる。この辺が面白いところではないかなというふうに思っています。

最後ちょっと駆け足になりましたけれども、多次元尺度構成法の基礎とその応用ということでお話しさせていただきました。どうもありがとうございました。

付記

本稿の内容は、日本認知心理学会・研究法研究部会 第1回研究会の発表内容に基づくものです。

文献

- Abelson, R. P. (1954-55). A technique and a model for multi-dimensional scaling. *Public Opinion Quarterly*, **Winter**, 405-418.
- 朝野熙彦.(2010). 最新マーケティング・サイエンスの基礎. 講談社
- 千野直仁 (2003). 多次元尺度構成法講義ノート
<http://www.aichi-gakuin.ac.jp/~chino/part-time/handai.pdf>
- 小杉考司 (2019). 言葉と数式で理解する多変量解析入門 北大路書房
- 小杉考司 (2019). 日本認知心理学・研究法研究部会・第一回 MDS の巻伴走サイト https://kosugiti.github.io/JSCP_MDS_2019/
- 岡太彬訓・今泉忠 (1994). パソコン多次元尺度構成法 共立出版
- 荘島 宏二郎 (2011) 非対称フォン・ミーゼス尺度法
<http://www.rd.dnc.ac.jp/~shojima/ams/jindex.htm>
- Stalans, L. J. (1994). Multidimensional scaling. In G. Grimm & P. R. Yamold (Eds.), *Reading and understanding multivariate statistics* (pp. 137-168). Washington DC: American Psychological Association. (Stalans, L. J. 高田 菜美・山根 嵩史 (訳) (2016). 多次元尺度構成法 L. G. グリム・P. R. ヤーノルド (編) 小杉 考司 (監訳) 研究論文を読み解くための多変量解析入門——重回帰分析からメタ分析まで—— (pp. 127-153) 北大路書房)
- 高根芳雄 (1980). 多次元尺度法 東京大学出版会
- 豊田秀樹(編)(2018). たのしいベイズモデリング, 北大路書房.
- 吉田戦車 (1994). 伝染るんです。(4) 小学館