

心理学データ解析 (心理学統計法)

担当: 小杉考司

Last Compiled on 2025.4.7

はじめに

授業のテーマ

「見えないもの」をみよう・測ろう・考えようとするのが心理統計の世界である。

一年を通じて伝えたいポイント

母数を知るための推測 大学での統計, 心理統計は標本の記述統計量を検証するのではなく, 母集団を表す数字を求めるための推測統計である。

確率を用いた推論 母数を知るための方法は3つある。代表値を用いたモーメント法, 確率分布をつかった最尤法, ベイズ法である。

要因計画と線形モデル 心理学では, 要因計画による研究条件をコントロールし, 平均値の差をもちいた考察 (平均因果効果) をおこなうことが多い。そしてそれらは総じて線形モデルとして表現できる。

モデル比較と意思決定 モデルによる母数の推定だけではなく, そこから一定の「結論」あるいは「意思決定」を行うための方法として, モデル比較や NHST といった方法がある。

統計環境 R による実践 統計環境 R を利用して, 理論的な理解だけでなく実際に計算ができることも身につけるべき技術である。

1 心理学と統計法

1.1 授業内容

1.1.1 概要

心理学を学ぶに当たって、なぜ統計学を学ばなければならないのかを考える。そのためには、心理学がサイエンスであること、サイエンスであるとは客観的であること、積み重ねが必要であることを知る必要があり、そのために数量的表現が有用であることを理解する。

1.1.2 コマ主題細目

心理学とはどういう学問か 心理学は初等教育の中には取り入れられていないにもかかわらず、心や気持ちといった用語は日常的に頻繁に用いられ、かつ重要視されているものでもある。「心理学の過去は長いが歴史は短い」といわれるが、どのようにして近代科学の仲間入りをしたのかを理解するため、サイエンスとしての枠組みを改めて問い直す。

→ 三浦 (2017) の P.1-7.

心理学の 2 つの流派 心理学には基礎と応用、という区別がされることがあるが、内容やその方法論的な区分から行けば実証主義的アプローチと理解主義的アプローチの 2 つに分類する方が良い。また基礎と応用という区分では、基礎が先に、下にあって応用が後に、上にくるような印象を持つが、それらは理学と工学の違いのように、併存して然るべきものである。

→ 下山 (2001) による研究方法の分類

近代科学の特徴 心理学は科学である、というならばその科学とはそもそもなんだろうか。ここでは近代科学の特徴として実験、数学的現象主義、機械論、要素還元主義をあげる。ニュートンやガリレオが行なったのは、関数による数学的現象主義であり、心理学そのものは基本的にその枠組みの中に位置する。すなわちアウトカム変数に対して心というモデルを立てて、そのモデルの正しさを検証することが 1 つの目的である。アウトカム変数は数量化されていなくてもよいが、数量化されていた方が正確性、客観性、比較可能性、要約可能性に優れており、かつ統計モデルの活用により表面的に現れる数字から見えない情報を引き出すことが可能になっている。

→ 三浦 (2017) の P.79-81. → 数値化の意味については川端・荘島 (2014) の P.1-6

1.1.3 キーワード

- 心理学における 2 つの検証方法
- 近代科学としての心理学
- 数学的現象主義

1.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 1. 心理学における実証的研究法 (量的研究及び質的研究); A 科学と実証

1.2.1 予習・復習課題

■予習 心理学入門の入門書として、道又 (2009) と中西・今田 (2020) などを一読しておくことで、これから4年間かけて学ぶ心理学についての事前のイメージを再構築すると良い。ほかにも、日本心理学会が一般向けに出版している「心理学ワールド」が読みやすい入門雑誌であろう。あるいは心理学についてのイメージ調査については菊池 (2018) が詳しく、心理学の歴史についての専門書としては大芦 (2016) なども参考にすると良い。

■復習 本学における心理学の学習が始まる前に持っていた、心理学に対する事前のイメージをいったんわすれ、まずは広く心理学の分野について眺めた上で、入学を志望したときに持っていた心理学についてのイメージを改めて位置づけてもらいたい。心理学のより細かい専門領域に対する位置付けを把握した上で、心理学統計法のような研究方法が統一され共有されていること、とくに初年時の必須科目として位置付けられていることの意義を自分なりに見出せるように、ノートなどに授業の感想と現段階での今後の展望を書き留めるなどすると良い。

2 心理学と計算機

2.1 授業内容

2.1.1 概要

心理統計の授業では、多くの数値を一気に扱うため計算機を用いる。Excel/Numbers/Calcs のような表計算ソフトでは不十分なため、R と呼ばれる統計専門のソフトウェアを用いる。R の使い方については次回以降に詳細が論じられるが、それに先立って計算機の基礎的な概念、操作方法を理解する。

2.1.2 コマ主題細目

電子計算機の基礎 計算機には 5 つの機能がある。すなわち、入力、出力、演算、制御、記憶装置とよばれるもので、対応するデバイスの組み合わせで計算機が構成される。また、BIOS、OS、アプリケーションの違いなど、計算機の内部構造についての理解は、人間をコンピュータとみなして考えるとき、メタファとして有用である。

情報の単位 電子計算機はすべて 0/1 の情報として扱う。さまざまな単位と、対応する装置ではどの程度の大きさを扱えるのかについて、基本的な知識を身につける。ビット、バイト、キロバイトからテラバイトまでの関係性や、テキストデータと画像・動画データの容量の違いなどを理解することで、扱うデータの規模感を把握する。

ファイルの種類と拡張子 コンピュータで扱うファイルには様々な種類があり、拡張子によって区別される。特に統計解析で扱うことの多いテキストファイル (.txt)、CSV 形式 (.csv)、Excel ファイル (.xlsx) の違いを理解し、それぞれの特徴と用途を把握する。また、R で扱うデータファイル (.RData) やスクリプトファイル (.R) についても概説する。

ファイルの場所 コンピュータ内でのファイル管理の仕組みとして、ディレクトリ (フォルダ) 階層構造について学ぶ。カレントディレクトリ、相対パス、絶対パスなどの概念を理解し、R を使った解析においてファイルの読み込みや保存をする際に必要となる基礎知識を習得する。クラウドストレージの利用についても触れ、バックアップの重要性を理解する。

キーボードとショートカット 効率的なコンピュータ操作のために、キーボードの配置やタイピングの基本、よく使うショートカットキー (コピー・ペースト、検索、保存など) について学ぶ。特に R や RStudio での作業に役立つキーボード操作を中心に解説し、実際に操作しながら習得していく。

2.1.3 キーワード

- コンピュータの基本構造 (入力・出力・制御・演算・記憶)
- ファイル形式とデータ管理
- 効率的な PC 操作

2.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 該当しません**2.2.1 予習・復習課題**

■予習 大学の PC ルームやメディアセンターで利用できるコンピュータの基本操作方法を確認しておくこと。自分のノート PC を授業に持参する場合は、事前に OS のバージョン、ストレージの空き容量、メモリ容量を確認しておくこと。また、普段使用しているデバイス（スマートフォンやタブレット）と比較して、PC ではどのような操作が異なるかについて考えてみると良い。

■復習 授業で紹介したファイル操作（作成・保存・移動・削除）を自分の PC で実践してみる。特に、新しいフォルダを作成し、その中にテキストファイルを作成・保存する練習をすること。また、キーボードショートカット（Ctrl+C, Ctrl+V, Ctrl+Z, Ctrl+S など）を意識的に使用する習慣をつけるよう心がけること。次の授業までに、R とそのインターフェースとなるソフトウェア (RStudio) をインストールしておくこと。

3 統計環境 R の導入

3.1 授業内容

3.1.1 概要

携帯電話やタブレットが誰でも身近に利用できるようになってきてはいるが、それらはあくまでも受動的にサービスを利用する媒体であり、執筆や計算などユーザが能動的に電子機器を利用するにあたっては、パソコンの利用は不可欠である。本講義だけでなく、心理学の研究をする上でも当然パソコンの利用は必須であり、とくに統計的データ処理をするにあたってはより深く知る必要がある。また、Excel や Numbers などの表計算でも記述統計量やグラフの作図はできるが、それ以上の「複数の変数を扱った多変量解析」や「群間の平均値の差を細かく検証する」といった心理統計的な作業には不十分なのである。そこで、より専門的なツールとして統計環境 R を導入する。統計環境 R は専修大学人間科学部心理学科が公式に採用している統計環境であり、これを活用する総合開発環境である RStudio の基本的な使い方、結果の再現性を担保するための工夫などについても言及する。それに先立って、コンピュータの基礎知識についても解説する。

3.1.2 コマ主題細目

統計環境 R/RStudio の準備 心理学科では基本的に統計環境 R を共通のツールとして用いる。統計環境 R はフリーソフトウェアであり、誰でも無償で利用できる環境である。また、R 単体では簡素な言語的やりとりができるに過ぎない。ファイルの操作や描画機能、記録、参照、管理など統計言語を扱うに当たって必要な総合的環境として RStudio を併用することで、その利便性は飛躍的に向上する。まずは R/RStudio を導入し、基本的な画面構成やプロジェクトによる管理ができることを確認する。

→ R についての入門書は多くあるがここでは小杉 (2019) の P.1-7 を挙げておく。同じく RStudio を使った入門書は多くあるがここでは小杉 (2019) の P.7-24 を挙げておく。

R の基本操作 環境が整ったところで、R の基本的な操作を見ていく。まずは R スクリプトをつかって、四則演算や関数計算、データの形式やパッケージの導入と利用について理解する。スクリプトペインとコンソールペインの違いを理解し、またスクリプトとして清書しておくことで計算手続きの記録がつけられることの利点を理解する。

基本的なデータ型と演算 R では Numeric(数値), Character(文字列), Logical(論理値) などの基本的なデータ型があり、それぞれの特徴を理解する。またベクトル、行列、データフレームといったデータ構造の違いを学び、基本的な統計量の算出方法を実践する。

→ 小杉 (2019) の P.25-40

3.1.3 キーワード

- R/RStudio
- スクリプトとコンソール
- データ型とデータ構造
- パッケージ管理

3.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; 2. 統計に関する基礎的な知識; C エクセル, R, SPSS 入門 (1) 統計分析のための言語の手ほどき プログラムの基本的操作

3.2.1 予習・復習課題

■予習 事前に R(<https://www.r-project.org/>) と RStudio(<https://posit.co/download/rstudio-desktop/>) をダウンロードしてインストールしておくこと。インストールに問題がある場合は、授業開始前に教員や TA に相談すること。

■復習 授業で学んだ R/RStudio の基本操作 (四則演算, 関数の使用, データの入力など) を実際に自分の PC で繰り返し練習すること。特に, スクリプトを用いた操作に慣れること。また, 簡単な数値データを用いて平均値や標準偏差を算出する練習を行うこと。

4 数値データの種類

4.1 授業内容

4.1.1 科目の中でのこのコマの位置づけ

前回までに R の基本的な操作方法を学び、コンピュータを用いた統計解析の基礎を理解した。本コマではデータの特性について学ぶことで、適切な統計的手法を選択するための基盤を形成する。

4.1.2 概要

目に見えない心という対象を研究するために、言語や行動によって得られた反応を記録しなければならない、記録されたものは集計、分析のために数値化される。数字やその統計的分析結果が、対象を必ずしも反映していないという直観に反するかもしれないが、数字の限界は言語の限界と同じであり、言い換えれば論述・論考・論証の限界に等しいすべての科学的領域をカバーする。ここでは数値化することで何が得られるか、その長所を理解するとともに、データの特性に応じた適切な分析方法を選択するための基礎知識を身につける。

4.1.3 コマ主題細目

データと変数 個人の特徴を記述するにあたって、心理統計ではすべての情報を数値化する。これらの数字は人によって変わりうるので、一般に変数と呼ぶ。また、目に見える変数を観測変数やアウトカムとよぶが、目に見えない心の状態をデータから取り出したりする場合は、それを潜在変数と呼んで区別する。変数の性質を理解することは、適切な統計的手法を選択する上で重要である。

→ データの分類については、山田・村井 (2004) の P.18-21。

コンピュータにおけるデータ型 統計ソフトウェアでは、データの特性に応じて異なる型を使用する。バイナリ (二値) データは 0/1 の二つの値のみを取る変数であり、性別や有無といった情報がこれに該当する。カテゴリカルデータは複数のカテゴリーに分類される変数で、数値的な順序を持たない。オーディナルデータは順序を持つカテゴリカルデータであり、アンケートの「全く同意しない」から「強く同意する」などの尺度がこれに相当する。連続変数は身長や体重などの無限の値をとりうるデータであり、コンピュータ上では離散的に扱われるが理論上は連続的である。それぞれのデータ型は R などの統計ソフトウェアで異なる方法で処理され、適切なデータ型の指定が分析の前提となる。

→ R におけるデータ型については馬場 (2020) の P.65-74 を参照。

尺度水準による分類 すべての数字が四則演算が可能かと言われると、そうではない。尺度水準とは、数値が持つ情報や性質の段階を表す言葉であり、Steavens, S. が提唱した 4 段階の区別がとくに重要である。名義尺度 (nominal scale) は単なる分類を表し、順序尺度 (ordinal scale) は順序関係を表す。間隔尺度 (interval scale) は等間隔性を持ち、比率尺度 (ratio scale) は絶対的なゼロ点を持つという特徴がある。これらの尺度水準の違いは、適用可能な統計的手法を決定するため、得られた数字データがどの水準に該当するかを判別することができるようにならなければ、誤った統計処理をすることになるため注意が必要である。

→ 川端・荘島 (2014) の P.9-16, あるいは山田・村井 (2004) の P.22-25。

コンピュータデータ型と尺度水準の関係 コンピュータにおけるデータ型と心理統計学における尺度水準は緊密に関連している。例えば、名義尺度はカテゴリカルデータとして、順序尺度はオーディナルデータとして扱われる。間隔尺度と比率尺度は通常、連続データとして扱われるが、その解釈と扱いには重要な違いがある。R などの統計ソフトウェアでは、これらの違いを適切に指定することで、データの特徴に合った分析が可能となる。実際の研究においては、測定方法や理論的背景を考慮して、データの性質を適切に判断することが求められる。

4.1.4 キーワード

- 観測変数と潜在変数
- バイナリ・カテゴリカル・オーディナル・連続データ
- 確率分布とデータ型
- 尺度の四水準 (名義・順序・間隔・比率)

4.2 授業情報

■コマの展開方法 講義と実習

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 1. 心理学における実証的研究法 (量的研究及び質的研究); C 実証の手続き

4.2.1 予習・復習課題

■予習 一般に、データに基づいて (Evidence Based) 論証・考察することが主観的な思い込みから脱却する第一歩である。尺度水準については心理学や心理統計の入門的テキストでは必ず触れられる項目であり、川端・荘島 (2014), 小杉 (2018, 2019), 三浦 (2017), および山田・村井 (2004) などを手にとって確認しておくが良い。また、R を用いたデータ操作の基本について、前回の復習を行っておくこと。

■復習 身の回りにあるさまざまな数字に注意を払い、それらがどの尺度水準の情報を持っているか、またコンピュータ上ではどのようなデータ型で扱うべきかを考えることで理解を深める。例えば、学生証番号、授業の成績評価 (S/A/B/C/D)、テストの点数、体重などを考え、それぞれがどの尺度水準に該当するか、R ではどのデータ型で表現すべきかを整理してみる。また、各種の心理尺度 (例: リッカート尺度) がどの尺度水準に該当するかについても考察し、適切な統計的処理について考えてみると良い。次回の授業に向けて、R で実際にさまざまなデータ型を扱う練習をしておくこと。

5 相関係数と R による実践

5.1 授業内容

5.1.1 概要

これまでは 1 変数の記述統計について学んできたが、本コマからは複数の変数間の関係性について考察する。変数同士の関係性を数値化して表す指標として、共分散や相関係数について学ぶ。また、変数の尺度水準に応じて適切な関連指標が異なることを理解し、それぞれの特徴と適用条件について理解する。相関係数は直線関係についての指数であることに注意し、散布図による可視化の重要性を理解する。実際のデータを用いて R による計算と可視化の演習を行うことで、相関分析の実践的技能を身につける。

5.1.2 コマ主題細目

二変数間関係の指標 変数同士の関係を記述する方法として、クロス集計表、共分散、相関係数などがある。クロス集計表はカテゴリカルデータの関連を表現する基本的な方法である。共分散は二つの変数がどのように連動して変化するかを数値化したものだが、変数の単位に依存するという欠点がある。そこで単位の影響を取り除くために標準化したものが相関係数であり、-1 から +1 の範囲で二変数間の線形関係の強さを表す。これらの基本的な考え方と計算方法について理解する。

→ 山内 (2010) の P.62–68, 山田・村井 (2004) の P.44–50, 川端・荘島 (2014) の P.48–52.

尺度水準に応じた相関係数 変数の尺度水準によって適切な関連指標は異なる。名義尺度の関連はカイ 2 乗値やファイ係数、クラメールの V 係数などで表される。二値データ間の関連を表すテトラコリック相関、順序尺度間の関連を表すポリコリック相関、順序尺度と連続変数の関連を表すポリシリアル相関などがある。連続変数間の関連を表す指標としては、ピアソンの積率相関係数が最も一般的であるが、線形関係以外も検出できるスピアマンの順位相関係数やケンドールのタウ、さらに最近では非線形関係も捉えられる MIC(Maximum Information Coefficient) なども開発されている。これらの指標の特徴と適用条件を理解し、データの性質に応じた適切な指標を選択できるようになる。

→ 様々な相関係数について川端・荘島 (2014) の P.53–58, 山田・村井 (2004) の P.51–59, MIC やテトラコリック相関などカテゴリカルな相関係数については Shojima (2022) の 2 章を参照。

R による相関分析の実践 相関係数の計算と検定は R を用いて簡単に行うことができる。cor() 関数を用いた相関係数の計算方法について学ぶ。また、polycor パッケージを用いたポリコリック相関やポリシリアル相関の計算方法、minerva パッケージを用いた MIC の計算方法についても触れる。実データを用いた演習を通じて、相関分析の実践的な技術を身につける。

→ R による相関分析については小杉 (2019) の P.110–120 を参照。

相関の落とし穴と可視化の重要性 相関係数だけを見て判断することの危険性について理解する。同じ相関係数の値でも、散布図のパターンは大きく異なる可能性がある。この点を示すための有名な例として、アンスコム の例や、より現代的な datasauRus パッケージに含まれる Dinosaur dataset がある。これらの例を通じて、相関分析を行う際には必ず散布図を描いてデータの分布を確認することの

重要性を学ぶ。また、外れ値の影響や非線形関係、層別の問題など、相関分析における様々な落とし穴について理解し、それらを回避するための方法について学ぶ。

→ `datasauRus` パッケージには、記述統計が同じだがプロットがまったく異なる 13 のケースが含まれていて参考になる。可視化の重要性については Healy (2018 瓜生他訳 2021) を参照。

5.1.3 キーワード

- 共分散と相関係数
- カイ 2 乗値, ファイ係数, テトラコリック相関, ポリコリック相関
- ピアソンの積率相関係数, スピアマンの順位相関, MIC
- 散布図による可視化

5.2 授業情報

■コマの展開方法 講義と演習

■標準シラバスにおける位置づけ 科目番号 5 心理学統計法; 1. 心理学で用いられる統計手法; B 記述統計: 相関

5.2.1 予習・復習課題

■予習 分散, 標準偏差, 標準化の手続きについて再確認しておくこと。分散は自分自身との共分散という形で理解することができるし, 共分散を標準化したものが相関係数になるからである。また, R の基本操作とデータ読み込みの方法について復習しておくこと。

■復習 授業で扱ったサンプルデータを用いて, 以下の課題を行い, Rmarkdown ファイルにまとめて提出すること。1. 複数の変数間の相関係数を計算し, 相関行列を作成する 2. 相関係数のヒートマップを作成する 3. 散布図を描き, 相関係数の値と視覚的な関係性を比較考察する 4. `datasauRus` パッケージの例を再現し, 同じ相関係数でも異なるパターンが生じることを確認する 5. 身近な例から二つの変数を選び, それらの関係を相関分析によって検討する

また, 相関関係と因果関係の違いについて理解を深めるため, 日常生活や報道などで見られる「相関関係を因果関係と誤解している例」を一つ挙げ, なぜそれが誤りであるかを説明すること。これはデータリテラシーの問題でもあり, 統計的思考の基本でもある。

6 回帰分析の基礎

6.1 授業内容

6.1.1 概要

前回学んだ相関係数は変数間の関連性を示すものであるが、一方の変数から他方の変数を予測・説明する枠組みとして回帰分析がある。本コマでは最も基本的な単回帰分析について学び、予測モデルの構築と評価の方法を理解する。また、相関関係と因果関係の違いについて考察し、因果推論の基礎的な考え方を学ぶ。古典的テスト理論から測定に含まれる誤差を知り、誤差が正規分布に従っていると考え、平均の位置を推定することができるようになった考え方を拡張し、他の変数との関係から平均の位置が関数で表現される最も単純な場合である線形モデルについて考える。

6.1.2 コマ主題細目

予測と線形モデル 変数間関係を記述する最も単純な形として、一次関数を用いたものを考える。現象に対して関数で記述する数学的現象主義、変数を説明変数と被説明変数として捉えること、予測に誤差が加わって現象となることなど、これまでの議論をデータと数式に置き換えただけであることをしっかりと理解する。さらに単回帰分析の数式的記述について理解する。とくにデータ (X_i, Y_i) 、予測値 $(\hat{Y})$ 、誤差 (e_i) 、回帰係数 (b_0, b_1) として表される記号の細部にまで注意を払い、その意味を理解する。

→ 豊田 (2017) の P.27-28.

最小二乗法 最小二乗法によって推定値が算出できることを理解する。数式の細かい展開はここでは行わず、最小化する目的関数の直観的理解に止める。また算出された係数 (傾きと切片) がどのような意味を持つのか改めて理解する。傾きは説明変数が 1 単位変化したときの被説明変数の変化量を表し、切片は説明変数がゼロのときの被説明変数の値を表す。これらの係数の解釈と実践的な意味について考察する。

→ 豊田 (2017) の P.30

残差と決定係数 回帰係数が計算されれば、実際のデータに対して予測値を計算することができ、予測がどの程度適合していたかを検証することができる。予測値と被説明変数との差分を残差 (residuals) と呼ぶが、この両者の相関係数を重相関係数、その二乗したものを決定係数 (R^2) と呼んでモデルの適合度を測る指標にすることを理解する。決定係数は被説明変数の分散のうち、モデルによって説明される割合を表し、モデルの説明力を示す重要な指標である。

→ 野島他 (2019) の P.58-60.

相関関係と因果関係 相関関係があるところに、因果的な説明をしてしまいがちであるが、因果的な説明をするためには相関以上に注意が必要である。因果関係を見るために必要な、時間的先行性、他の原因がないこと、他の結果に至らないことなどの諸条件を確認する。回帰分析は変数間の関係を数理モデルとして表現する強力なツールであるが、それだけでは因果関係を示すには不十分であることを理解する。回帰係数が因果効果として解釈できる条件と、そうでない場合の違いについて理解する。

→ 相関と因果の違いについては、豊田他 (1992) に簡潔にまとめられている。

6.1.3 キーワード

- 一次関数と線形モデル
- 予測値, 誤差, 回帰係数
- 最小二乗法と決定係数
- 相関関係と因果関係

6.2 授業情報

■コマの展開方法 講義と演習

■標準シラバスにおける位置づけ 科目番号 5 心理学統計法; 1. 心理学で用いられる統計手法; C 記述統計: 回帰

6.2.1 予習・復習課題

■予習 前回学習した相関係数の概念と計算方法について復習しておくこと。また、一次関数の基本的な性質 (傾きと切片の意味) についても確認しておくが良い。

■復習 山内 (2010) の P.62–81 は相関係数と回帰分析を 1 つの章にまとめて解説しているので参考にすると良い。最小二乗法による回帰係数の算出については、偏微分連立方程式を用いれば簡単に計算できるが、偏微分連立方程式とは何かに遡って解説している本として小杉 (2018) の P.104–112 がある。

また、授業で学んだ内容を応用して以下の課題に取り組むこと:

1. 実際のデータセットを用いて単回帰分析を実施し、回帰係数と決定係数を算出する
2. 得られた回帰式を用いて予測を行い、残差をプロットする
3. 身近な例から「相関関係と因果関係の混同」の事例を見つけ、なぜそれが因果関係とは言えないのかを説明する
4. 因果関係を主張するためには、相関関係に加えてどのような条件や証拠が必要かを考察する

7 重回帰分析と変数間の影響

7.1 授業内容

7.1.1 概要

今回は単回帰分析について学んだ。今回は複数の説明変数を用いる重回帰分析へと発展させる。回帰分析によって、データにモデルを当てはめて係数の推定値を得るということを行ったが、そのこととモデルが適していたかどうかは別である。この点を理解するためには、モデルの適合度について考える必要がある。モデルが適合しているということを踏まえた上で、説明する変数が増える重回帰分析へと議論は展開する。ここでは説明変数の読み取り方への注意を促すべく、偏相関係数の考え方を解説し、偏回帰係数、標準化偏回帰係数などの用語について理解する。

7.1.2 コマ主題細目

重回帰分析の基本 説明変数が複数に増えた場合の回帰分析は、重回帰分析 (Multiple Regression Analysis) と呼ぶ。数式的な表現としては項が増えるだけであるが、データのプロットを考えると三次元以上の回帰平面に拡張されていること、それでも面は湾曲していない、すなわち線形関係であることを理解する。単回帰分析と同様に、最小二乗法によって回帰係数を推定するが、行列計算を用いることで効率的に解が求められることにも触れる。

→ 小杉 (2019) の P.168, 図 10.7 の回帰平面図を参照

交絡と統計的制御 複数の変数が相互に関連し合っている場合、見かけ上の関連 (擬似相関) が生じることがある。第三の変数によって見かけ上の相関が生じている場合、その変数を交絡要因 (confounder) と呼ぶ。重回帰分析では、これらの交絡要因をモデルに含めることで、その影響を統計的に制御することができる。これは実験的に条件を必ずしも統制できない観察研究において、統計モデル的に統制する手法として重要である。交絡要因の影響を取り除いた上での変数間関係を評価することができる。

→ 交絡に関しては豊田他 (1992) の P.45-58 を参照。

部分相関と偏相関 重回帰分析における係数の解釈のために、部分相関と偏相関の概念を理解する。回帰分析によって、被説明変数の分散が説明される部分と残差の部分に分離させられること、残差と説明変数は相関しないことを確認し、残差は説明変数の影響を除外したものと考えることができる。これは条件を必ずしも統制できない事象において、統計モデル的に統制していることでもある。偏相関係数は他の変数の影響を取り除いた 2 変数間の純粋な関連の強さを表す指標であり、変数間の直接的な関係を評価する上で重要な概念である。

→ 川端・荘島 (2014) の P.59-65, 山田・村井 (2004) の P.60-65。

偏回帰係数と解釈 部分相関の説明において、残差が分離されることを見てきたが、残差同士の相関関係はとくに偏相関ということができる。これは共通する変数からの影響を除いた変数間関係であり、これを見るだけでも変数間関係を見ることにつながる。とくに残差変数を使っている回帰分析を行うとき、得られる係数が偏回帰係数である。重回帰分析における回帰係数はこのように条件付きで考えなければ

ばならない。偏回帰係数は「他の説明変数の値を一定に保った場合に、その説明変数が1単位変化したときの被説明変数の変化量」として解釈できることを理解する。また変数を増やすことで重相関係数が増加することも理解する。

→ 小杉 (2018) の P.60–63.

標準化係数と多重共線性 回帰係数については、素点による回帰係数が理解しやすいが、単位に左右されるものでもあるので相対比較をする場合はすべての変数を標準化した標準化解を用いた方がわかりやすい。標準化解を用いても適合度は変化しないことなど留意点とともに有用性を理解する。また、複数の説明変数間に強い相関がある場合、多重共線性 (multicollinearity) の問題が生じる可能性があることを理解する。多重共線性があると、回帰係数の推定値が不安定になり、解釈が困難になるという問題が発生する。VIF (Variance Inflation Factor) などの診断指標や、多重共線性が検出された場合の対処法についても簡単に触れる。

→ 小杉 (2018) の P.66–67, 多重共線性については清水・荘島 (2017) の P.86–89 を参照。

7.1.3 キーワード

- 重回帰分析と回帰平面
- 交絡要因と統計的制御
- 部分相関と偏相関
- 偏回帰係数と標準化偏回帰係数
- 多重共線性

7.2 授業情報

■コマの展開方法 講義と演習

■標準シラバスにおける位置づけ 科目番号 5 心理学統計法; 1. 心理学で用いられる統計手法; C 記述統計: 回帰

7.2.1 予習・復習課題

■予習 単回帰分析について、数式的表現や最小二乗法の意味について概念的な理解で十分なので再確認しておくこと。重回帰分析では説明変数が複数になるため、解析的算出には行列計算の知識が必要になる。行列計算については、二年次の統計に関する講義で詳しく扱う予定であるが、興味があるものは「線形代数」をキーワードに予習すると良い。初学者には結城 (2018) または岡太 (2008) が適している。

■復習 豊田 (2017) の P.89–98 は本時の内容がわかりやすくまとめられているので、通読しておくことと良い復習になる。また、授業で学んだ内容を応用して以下の課題に取り組むこと:

1. 提供されたデータセットを用いて、R で重回帰分析を実施する
2. 各説明変数の偏回帰係数と標準化偏回帰係数を求め、その意味を解釈する
3. モデルの決定係数と調整済み決定係数を算出し、モデルの適合度を評価する

4. 多重共線性の診断を行い, 問題がある場合は対処法を考察する
5. 交絡要因の影響を統制した結果, 単回帰分析と比較してどのような変化があったかを考察する

重回帰分析は心理学研究で広く用いられる手法なので, Grimm・Yarnold (1994 小杉他訳 2016) の第二章などを読んで, どのような使い方をされているのか, また前提となっている条件や利用上の問題点などについても理解を深めておくこと。

8 Rをつかった回帰分析

8.1 授業内容

8.1.1 概要

回帰分析や重回帰分析について、基本的な概念や用語・推定値の意味は前回までに習得済みであるが、それが統計環境では実際どのようにして算出されるか、算出された数字はどのような意味があるかについて具体的に理解することを目的とする。座学の中では記号としてしか意味を持たなかったものが、実際の数字として得られるとより理解が深まることにもなる。統計環境は R を用いるが、ほかの統計パッケージを用いても同様の出力が得られる。統計環境が変わる可能性もあるので、出力画面の位置で内容を覚えるのではなく、意味的な理解をしておく必要がある。

8.1.2 コマ主題細目

散布図の描画 線形回帰はその名の通り線形関係をモデルで表現するものであり、非線形な関係や外れ値の存在などがあれば適切な推定値が得られたことにはならない。分析をする前にまずデータの様子を確認するために、散布図の描画は必ず行わなければならない。R の基本的な plot 関数を用いた散布図の描画方法を学び、データの視覚的な確認方法を理解する。また、複数の変数がある場合の散布図行列 (pairs 関数) の活用方法についても理解する。

→ 小杉 (2019) の P.92–98.

線形回帰の実行 データが読み込まれれば、lm 関数によって最小二乗法による回帰分析を行うことができる。R による関数関係の記述 (formula) の方法を理解し、出力結果から適切な読み取りができるよう、どのような用語でどのような出力がなされているかをしっかり確認する。特に、切片 (Intercept) と傾き (係数) の推定値、標準誤差、t 値、p 値の意味と解釈について理解する。

→ 小杉 (2018) の P.54–60.

回帰診断と Q-Q プロット lm 関数によって得られた回帰モデルの適切性を診断することが重要である。回帰分析の結果オブジェクトを plot 関数に渡すことで、残差プロット、Q-Q プロット、スケールロケーションプロットなど、モデル診断に有用なグラフが得られる。特に Q-Q プロットは残差の正規性を確認するための重要なツールであり、理論上の正規分布と実際の残差の分布を比較することで、回帰モデルの前提条件が満たされているかを確認できる。残差プロットでは残差のパターン (漏斗状や系統的な曲線など) を確認し、モデルの妥当性や外れ値の影響を検討する方法を学ぶ。

→ 小杉 (2018) の P.68–72

残差と予測値の確認 lm 関数から返される結果オブジェクトの中には、予測値 $\hat{Y}$ が fitted.values、残差が residuals という変数名で保存されている。これらを取り出したり図示したりすることで、回帰分析の諸特徴を理解する一助となる。数理的な導出は小杉 (2018) の P.127–135 にあるが、具体的な数字としてこれを確認し、理解する。また、residuals.lm(), fitted.lm() などの関数を用いた残差や予測値の取り出し方も学ぶ。

→ 小杉 (2018) の P.54–60.

重回帰分析の実行 統計環境において重回帰分析を行うのは、説明項を 1 つ追加するだけで良いので作業としては容易い。ただし解釈に注意が必要であることを再確認し、また標準化偏回帰係数の算出方法についても理解を進める。また変数を増やす前と増やした後で、重相関係数が増加することを確認する。さらに、summary 関数から得られる調整済み決定係数 (Adjusted R-squared) や、AIC, BIC などのモデル選択指標についても理解する。

→ 小杉 (2018) の P.63–67.

8.1.3 キーワード

- R の基本 plot 関数
- lm 関数と formula 表記
- Q-Q プロットと残差診断
- fitted.value と residuals
- 標準化偏回帰係数

8.2 授業情報

■コマの展開方法 演習

■標準シラバスにおける位置づけ 科目番号 5 心理学統計法; 2 統計に関する基礎的な知識; C/D/E エクセル, R, SPSS 入門

8.2.1 予習・復習課題

■予習 第??講と同じ形式で行われる。R や RStudio, プロジェクトによる管理など基本的なことを改めて教示しないので、前回の復習をかねて空いている時間に PC ルームで基本的な R/RStudio の挙動について理解しておくが良い。また R の内部でどのように数字やデータが扱われるかについて、小杉 (2019) の P.24–36 などを参考にみておくが良い。さらに、単回帰分析と重回帰分析の概念的な理解を確認しておくこと。

■復習 授業で扱ったデータを用いて、以下の課題に取り組むこと:

1. 単回帰分析を実行し、結果を解釈する
2. 回帰診断プロット (特に Q-Q プロット) を描画し、モデルの妥当性を検討する
3. 重回帰分析を実行し、単回帰分析との結果の違いを比較する
4. 残差と予測値を取り出し、散布図上にプロットする
5. 標準化偏回帰係数を算出し、変数間の相対的な影響力を評価する

9 確率と統計的推論

9.1 授業内容

9.1.1 概要

母集団に確率分布を仮定し、そこから得られる標本分布の特徴を利用して母数を推定する方法について学ぶ。標本の平均値が母平均に一致することを確認し、また標本の分散は母分散に一致しないことを確認する。これまでは手元にあるデータを記述する方法を学んできたが、本コマからは未知の母集団の特性を推測するための統計的推論の基礎に入る。現実の研究では全数調査は困難であるため、標本から母集団へと推論することが必要となる。ここでの推定方法はモーメントを用いたものであるが、確率分布を用いる方法もあり、その場合の“確率モデル”の考え方が必要なことを学ぶ。

9.1.2 コマ主題細目

母集団と標本 推測統計学の基本的な概念である母集団と標本の関係を理解し、統計学的な推定（点推定、区間推定）の位置付けについて学ぶ。^{*1}用語としての母数（母平均、母分散、母比率など）と標本統計量（標本平均、標本分散、標本比率など）、および標本統計量の実現値といった用語の意味するところも慎重に理解しなければならない。特に母集団分布、標本分布、サンプリング分布の違いについて理解することが重要である。

→ 山田・村井（2004）の P.68–73

母集団分布とモデル ここでは母集団に正規分布を仮定する“モデル”について考える。統計的な推測は基本的に不良設定問題であり、何らかの前提をおかないと解くことができない問題がほとんどである。心理学の場合は母集団分布を正規分布とする仮定が選ばれることが多く、この場合は標本統計量の分布も計算しやすくなる。現実の分布が必ずしも正規分布に従わない場合でも、中心極限定理により標本平均は正規分布に近似できることが、この仮定の有用性を支えている。

→ 山田・村井（2004）の P.74–79.

正規分布の別の意味 正規分布はここまで、誤差の分布として解説してきたが、複数の要因が積み重なった時の全体的傾向として理解することもできる。平均値が集団の代表値としての意味もあること、また中心極限定理によって分布によらずその標本平均値が正規分布に従うことが示される。中心極限定理は、母集団分布がどのような形であっても、標本サイズが十分大きければ標本平均の分布は正規分布に近づくという重要な性質であり、推測統計学の基礎となる定理である。

中心極限定理については → 皆本（2015）の P.180–181.

標本分布と標準誤差 母集団が正規分布に従うとすると、標本統計量も正規分布に従うことが導出できる。ここでは標本統計量としての平均値を例にあげ、「標本統計量の分布」のことを標本分布と呼ぶこと、また標本分布の標準偏差をとくに標準誤差と呼ぶことを理解する。標準誤差は標本平均のばらつきを表す指標であり、サンプルサイズの平方根に反比例することを理解する。たとえばサンプルサイズが大

^{*1} 本書では検定をモデル比較の一種として扱うので、ここでは言及しない。

きくなったら、小さくなったらどのような値が得られるのかについて、具体的な例をみながら考えると理解が進む。

→ 標本分布については山田・村井 (2004) の P.90-97.

9.1.3 キーワード

- 母集団と標本
- 標本分布と標本統計量
- 中心極限定理
- 標準誤差

9.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; (2) 統計に関する基礎的な知識; A 統計分析の基礎: 母集団と標本

9.2.1 予習・復習課題

■予習 確率と確率分布の基本概念を復習しておくこと。特に標準正規分布の特徴とその利用方法について確認しておくことで理解が深まる。また、ここでは山田・村井 (2004) の P.68-97 が丁寧に解説されているので、事前に一読しておくことを強く勧める。

■復習 いよいよ記述統計をこえ、直接知り得ないものについて推論する段階に進む。基本的に解けない問題を解こうとするようなものであり、仮定やモデル、みなしている箇所について自覚的でなければならない。今回の講義の中で、何が仮定され、何が理論的に導出されたものなのかを区別できるかどうか、ノートを整理しておくことで良いだろう。また推測統計学の基本的概念である、母集団と標本の関係と推定の手続きについては、山田・村井 (2004) のテキストが最も丁寧に適切な記述をしているので、是非該当箇所を通読し、概念や用語を間違えないように注意されたし。

以下の点について理解を深めておくことも有効である:

1. 標本サイズと標準誤差の関係
2. 中心極限定理の意味とその応用
3. 母集団分布を正規分布と仮定することの意義と限界
4. 点推定と区間推定の違いと用途

10 統計的推定の基礎

10.1 授業内容

10.1.1 概要

これまでの授業では母集団分布と標本分布について学んできた。本コマでは標本から母集団の特性値(母数)を推定する方法について学ぶ。まず不偏推定量の概念を理解し、点推定と区間推定の違いについて学ぶ。標本統計量を用いて母数を推定する際の基本的な考え方と、推定値の精度や信頼性を評価する方法を理解する。さらに、推定と検定の関係についても触れ、統計的推論の基本的な枠組みを身につける。また、標準正規分布が統計的推論において果たす重要な役割についても理解を深める。

10.1.2 コマ主題細目

不偏推定量 母集団に正規分布を仮定し、標本平均の平均(期待値)が母平均に一致することを確認する。

このことから、標本統計量が推定値として利用できることを確認する。ただし、分散の場合は一致しない(不偏性がない)ことをみる。分散については不偏分散を考えることでこの不一致を解決することができる。標本統計量が母数の良い推定量であるための条件(不偏性、一致性、効率性)についても理解する。また、推定量と推定値の違いについても理解する。

→ 不偏推定量については山田・村井(2004)のP.98–103に詳しい。標本統計量のこれらの特徴については、小杉他(2023)のP.119–141も参照すると良い。

点推定と標準誤差 点推定とは、標本統計量を用いて母数の値を単一の値として推定する方法である。例えば、標本平均は母平均の点推定値となる。しかし、この推定値にはサンプリング誤差が含まれるため、その精度を評価する必要がある。標準誤差は推定値の精度を表す重要な指標であり、標本平均の標準誤差は母標準偏差をサンプルサイズの平方根で割ったものとして定義される。標準誤差の意味と解釈、およびサンプルサイズとの関係について理解する。また、中心極限定理により、サンプルサイズが大きくなると標本平均の分布が正規分布に近づくことも確認する。

→ 川端・荘島(2014)のP.107–110、山田・村井(2004)のP.116–119。

区間推定と信頼区間 標本平均は母平均の不偏推定量ではあるが、ある値をそのまま母平均であると考えるのは(標準誤差の考え方があるとは言え)やや断定的である。これに対して一定の幅を使って予測する区間推定という方法がある。推定に確率の言葉が入ってくるので不慣れな表現が出てくるが、注意深く解釈する必要がある。95

→ 川端・荘島(2014)のP.111–116、山田・村井(2004)のP.120–127。

標準正規分布と統計的推論 標準正規分布は統計的推論において中心的な役割を果たす。標本統計量を標準化することで標準正規分布に対応させ、その確率的な性質を利用して推定や後の検定を行うことができる。標準正規分布表の使い方や、標準正規分布に基づく信頼区間の構成方法について学ぶ。また、標準正規分布を用いた統計的推論の基本的な考え方と、これが後の検定につながることを理解する。

→ 山田・村井 (2004) の P.86–89, 川端・荘島 (2014) の P.36–44.

推定と検定の関係 推定は母数の値を推し量る作業であるのに対し、検定はある仮説の下での統計量の期待値と実際の統計量を比較することで、仮説の妥当性を判断する作業である。両者は密接に関連しており、推定は「どれくらいか」という問いに、検定は「差があるか否か」という問いに答えるものである。推定と検定の関係を理解することで、次回学ぶ帰無仮説検定の基礎を形成する。特に、信頼区間と検定の関係について理解を深める。

→ 川端・荘島 (2014) の P.81–85, 山田・村井 (2004) の P.104–110.

10.1.3 キーワード

- 不偏推定量
- 点推定と標準誤差
- 区間推定と信頼区間
- 標準正規分布
- 推定と検定の関係

10.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; F 推測統計: その考え方

10.2.1 予習・復習課題

■予習 前回までに学んだ標本分布, 標準誤差, 標準正規分布についての理解を確認しておくこと。特に, 標本統計量の標準化の方法と標準正規分布表の読み方について復習しておくことと理解が深まる。

■復習 授業で学んだ内容を定着させるために, 以下の点について理解を深めておくこと:

1. 不偏推定量の概念と, 標本分散と不偏分散の違い
2. 点推定と標準誤差の関係, およびサンプルサイズが推定精度に与える影響
3. 信頼区間の正しい解釈と, その構成方法
4. 標準正規分布を用いた統計的推論の基本的な考え方
5. 推定と検定の関係, および信頼区間と検定の関連性

統計的推定は, 実証的な心理学研究の基盤をなす重要な概念である。特に, 点推定だけでなく区間推定の考え方を理解することで, データの持つ不確実性を適切に評価できるようになる。統計的推定についての理解を深めるには, 川端・荘島 (2014) の P.104–116 や山田・村井 (2004) の P.98–127などを参考にするとうまい。

11 帰無仮説検定の基本

11.1 授業内容

11.1.1 概要

今回は統計的推定について学んだ。本コマでは、推定を発展させた統計的意思決定法である帰無仮説検定について学ぶ。帰無仮説検定は心理学研究において広く用いられている統計的推論の方法であり、データに基づいて仮説の妥当性を評価するための枠組みを提供する。本コマでは、帰無仮説検定の基本的な考え方と手順、さまざまな検定法（相関係数の検定、カイ二乗検定など）、そして有意水準や片側・両側検定の概念について理解を深める。帰無仮説検定は背理法を用いて効果の有無を判断する方法であり、結果の報告に関する表記法も含めてその作法を学ぶ必要がある。

11.1.2 コマ主題細目

統計的帰無仮説検定とその手順 統計的帰無仮説検定 (Null Hypothesis Significance Test; NHST) は帰無仮説に基づいてモデル比較をする一連の手続きである。前回学んだ推定と異なり、検定では特定の仮説 (帰無仮説) が正しいと仮定した場合に観測されたデータがどれだけ起こりにくいかを評価する。帰無仮説検定の基本的な考え方、帰無仮説と対立仮説の設定、検定統計量の算出、p 値の計算、判断基準の設定という一連の流れを学ぶ。また、検定統計量は標本統計量を標準化したものであり、統計量によって適切な分布 (標準正規分布、t 分布、カイ二乗分布、F 分布など) が対応することを理解する。

→ 手続きの一般化としては山田・村井 (2004) の P.108–109, 川端・荘島 (2014) の P.81–85.

相関係数の検定 相関係数の統計的検定は、母相関係数がゼロであるという帰無仮説を検証するものである。2 変数が無相関 (母相関係数 $\rho = 0$) という帰無仮説の下では、サンプルサイズ n と標本相関係数 r から計算される検定統計量が自由度 $n - 2$ の t 分布に従うことを理解する。この検定統計量の計算方法と、検定結果の解釈について学ぶ。また、相関係数の検定と同時に、相関係数の信頼区間を求める方法についても触れる。さらに、相関係数の大きさと統計的有意性の違いについても理解を深める。無相関検定は R では `'cor.test()'` 関数を用いて実行できる。

→ 川端・荘島 (2014) の P.98–103, 小杉 (2018) の P.58–60, 清水 (2021) の P.150–162.

カイ二乗検定 カイ二乗検定は、カテゴリカルデータの分析に用いられる検定法である。観測度数と期待度数の差を評価し、分布の適合度や変数間の独立性を検証する。カイ二乗検定には適合度検定と独立性の検定があり、それぞれの適用場面と計算方法について学ぶ。カイ二乗統計量の算出方法と、カイ二乗分布を用いた確率評価の方法について理解する。また、カイ二乗検定の適用条件 (期待度数が 5 以上など) や結果の解釈についても学ぶ。

→ カイ二乗検定については川端・荘島 (2014) の P.145–153, 山田・村井 (2004) の P.177–183 を参照。

有意水準 モデル比較という観点からは、どうしても「判定基準」を考える必要がある。NHST の文脈では、

それは希少性の指標である有意水準になる。有意水準の定め方は領域ごとに定める任意であり、心理学では一般に 5% が用いられている。このことは、複数のモデルの優劣を判定する数字であって、モデルの強さ、正しさを表現する数字ではないことに注意する。有意水準と p 値の関係、および有意水準の設定の仕方について理解する。また、p 値の正しい解釈についても学び、p 値の誤解や誤用についても理解を深める。

→ 山田・村井 (2004) の P.112–115, 川端・荘島 (2014) の P.86–89.

片側検定と両側検定 統計的検定には片側検定と両側検定があり、研究の仮説や関心に応じて適切な方法を選択する必要がある。両側検定は母数が特定の値と異なるかどうかを検証するのに対し、片側検定は母数が特定の値よりも大きい (または小さい) かどうかを検証する。それぞれの場合の帰無仮説と対立仮説の設定方法、臨界値の違い、p 値の計算方法、および結果の解釈について理解する。また、検定の方向性を事前に決定することの重要性と、結果を見てから検定方法を変更することの問題点についても学ぶ。

→ 山田・村井 (2004) の P.116–119, 川端・荘島 (2014) の P.90–93.

11.1.3 キーワード

- 帰無仮説検定と検定統計量
- 相関係数の検定 (無相関検定)
- カイ二乗検定
- 有意水準と p 値
- 片側検定と両側検定

11.2 授業情報

■コマの展開方法 適宜投影資料を用いる

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 2. データを用いた実証的な思考方法; C データの統計的記述
科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; F 推測統計: その考え方

11.2.1 予習・復習課題

■予習 前回学んだ不偏推定量, 点推定と区間推定, 標準正規分布の知識を復習しておくこと。特に, 標準誤差と信頼区間の概念, および推定と検定の関係について確認しておくことと理解が深まる。また, 相関係数の概念と計算方法についても再確認しておくこと。

■復習 授業で学んだ内容を定着させるために, 以下の点について理解を深めておくこと:

1. 帰無仮説検定の論理構造と基本的な手順
2. 相関係数の検定の手順と解釈 (無相関の検定, 検定統計量の計算, 結果の解釈)
3. カイ二乗検定の適用場面と使い方 (適合度検定と独立性の検定)

4. 有意水準と p 値の関係, および p 値の正しい解釈
5. 片側検定と両側検定の違いと, それぞれの適用場面

また, 実際に検定を行ってみることで理解が深まる。例えば, 相関係数の大きさ, 標本の大きさを任意の値にし, 無相関検定を実際にやってみると良い。さまざまな数値を入れることで, 検定結果がどのように変動するかを理解することは帰無仮説検定の本質を理解するのに役立つ。仮想的な数字では実感が得られにくい場合, 身の回りの相関がありそうなデータを具体例として検定し, 何が言えるのかを考えてみると良い。

12 帰無仮説検定の応用と注意点

12.1 授業内容

12.1.1 概要

前回までに帰無仮説検定の基本的な考え方と手順について学んだ。本コマでは、より実践的な検定法として母分散が未知の場合の平均値の検定 (t 検定) を学び、検定結果の報告方法や解釈について理解を深める。また、帰無仮説検定における 2 種類の誤り (第 1 種の誤りと第 2 種の誤り) について学び、検定力と効果量の概念を理解する。さらに、帰無仮説検定の限界や注意点についても考察する。特に、統計的有意性と実質的重要性の違い、検定の多重性の問題などについて理解し、帰無仮説検定を適切に活用するための視点を身につける。

12.1.2 コマ主題細目

1 つの平均値の検定 (t 検定) 前は相関係数の検定とカイ二乗検定について学んだが、本コマでは平均値の検定について学ぶ。実際の研究では母分散がわかっていることはほとんどなく、標本から不偏分散を用いて推定する必要がある。このような場合は標準正規分布ではなく、 t 分布と呼ばれる統計量が用いられる。 t 分布の特徴 (自由度によって形状が変わる、サンプルサイズが大きくなると標準正規分布に近づくなど)、 t 検定の手順、および t 分布表の使い方について学ぶ。また、1 サンプルの t 検定と、次回学ぶ 2 サンプルの t 検定の違いについても理解する。

→ 川端・荘島 (2014) の P.93–96, 山田・村井 (2004) の P.128–131.

検定における 2 種類の誤り 帰無仮説検定は確率的な判断を行う手法であるため、誤った判断を下す可能性が常に存在する。検定における誤りには 2 種類あり、第 1 種の誤り (Type I Error) は帰無仮説が正しいのに棄却してしまう誤り、第 2 種の誤り (Type II Error) は帰無仮説が誤りなのに棄却できない誤りである。有意水準 α は第 1 種の誤りを犯す確率の上限を表し、検定力 ($1 - \beta$) は第 2 種の誤りを犯さない確率を表す。検定力はサンプルサイズ、効果量、有意水準によって決まることを理解し、適切なサンプルサイズの決定方法 (検出力分析) についても学ぶ。

→ 山田・村井 (2004) の P.120–121, 検定の多重性については同書 P.158–161, 川端・荘島 (2014) の P.90–93 も参照。

効果量と実質的重要性 統計的有意性 (p 値が有意水準より小さいこと) は必ずしも実質的な重要性を意味するわけではない。特に大きなサンプルサイズでは、わずかな効果でも統計的に有意になりやすい。そこで、効果の大きさを標準化して表す効果量の概念が重要になる。代表的な効果量として、相関係数 (r)、Cohen's d 、 η^2 (イータ二乗) などがあり、それぞれの意味と解釈について学ぶ。また、効果量の大きさを評価する目安 (小・中・大) についても理解する。効果量は検定結果とともに報告することで、結果の実質的な意味を評価する助けとなる。

→ 効果量については大久保・岡田 (2012) の P.43–46, 川端・荘島 (2014) の P.97–98.

結果の報告と解釈 帰無仮説検定の結果を論文やレポートで報告する際の適切な方法について学ぶ。統計

量の値, 自由度, p 値, 効果量などを正確に報告することの重要性を理解する。また, p 値の捉え方として「帰無仮説が正しいという仮定のもとで, 観測されたデータやそれよりも極端なデータが得られる確率」であることを再確認し, p 値の誤った解釈 (帰無仮説が正しい確率, 効果の大きさなど) を避けるようにする。さらに, 統計的に有意であることの意味と限界について理解を深める。

→ 山田・村井 (2004) の P.132–133, 山内 (2010) の P.216–220.

帰無仮説検定の限界と注意点 帰無仮説検定は有用なツールであるが, いくつかの限界や注意点がある。

検定の多重性の問題 (多数の検定を行うと Type I Error の確率が増加する), 出版バイアス (有意な結果のみが報告される傾向), p-hacking (有意な結果が得られるまで分析を変更する) などの問題について理解する。また, これらの問題に対する対処法 (Bonferroni 補正, FDR 制御, 事前登録など) についても学ぶ。さらに, 近年の再現性危機と統計改革の動きについても触れ, 帰無仮説検定を補完する方法 (信頼区間, ベイズ統計, メタ分析など) についても概観する。

→ 検定の多重性については山田・村井 (2004) の P.158–161, 帰無仮説検定の限界については川端・荘島 (2014) の P.99–103.

12.1.3 キーワード

- t 検定と自由度
- 第 1 種の誤りと第 2 種の誤り
- 検定力と効果量
- 統計的有意性と実質的重要性
- 検定の多重性と帰無仮説検定の限界

12.2 授業情報

■コマの展開方法 適宜投影資料を用いる

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 2. データを用いた実証的な思考方法; C データの統計的記述

科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; F 推測統計: その考え方

12.2.1 予習・復習課題

■予習 前回学んだ帰無仮説検定の基本的な考え方や手順, 有意水準と p 値の関係について復習しておくこと。特に, 相関係数の検定やカイ二乗検定の手順と解釈について確認しておくことと理解が深まる。

■復習 授業で学んだ内容を定着させるために, 以下の点について理解を深めておくこと:

1. t 検定の手順と適用条件, および t 分布の特徴
2. 第 1 種の誤りと第 2 種の誤りの違い, および検定力との関係
3. 効果量の種類と解釈, および統計的有意性との違い
4. 帰無仮説検定の結果を適切に報告する方法

5. 帰無仮説検定の限界と注意点, および対処法

検定の話に入ると, 要因計画や線形モデルとどのようにこれらが接合するのが分かりにくい。回帰モデルの説明変数を離散化したのが要因計画→要因計画の効果の大きさ (傾きの大きさ) を質的に判断する手続きが, 平均値差の検定である, という関係であることを整理しておこう。推測統計学は, 基本的に母数の推定 (回帰係数のパラメータも母数である) をしているのであり, 推定量を何らかの基準で判断するのが検定, という繋がりである。もちろん検定を行わずに, 大きさをそのまま報告しても良い。むしろ昨今ではそちらの方が推奨されており, 些細な違いや「差がないこと」を「差があること」の根拠とすることの問題点が数多く指摘されている。帰無仮説検定の手続きは, 過去の論文に示された結果を読むための知識として重要であり, 本質である効果の大きさを推定していることを忘れないように。

13 前期のまとめと試験

13.1 授業内容

13.1.1 概要

これまでの14回の授業を通じて、心理統計の基礎となる概念と手法について学んできた。本コマでは前期の内容を総括するとともに、学習内容の理解度を確認するための試験を実施する。試験は授業時間内に行われ、これまでに学んだ記述統計、相関と回帰、確率分布、推測統計と検定などの内容から出題される。試験を通じて、心理統計の基本概念の理解と実践的な応用力が身についているかを確認する。

13.1.2 コマ主題細目

前期の総括 前期では「記述」と「因果と相関」を中心に学んできた。心理学における数量的アプローチの意義から始まり、データの可視化、記述統計量、相関と回帰分析、確率分布と統計的推論、そして帰無仮説検定まで、心理統計の基礎となる概念と手法を体系的に学習してきた。これらの知識は、後期の応用的な内容（分散分析、因子分析など）を学ぶための基盤となる。また、Rを用いた実習を通じて、統計解析の実践的スキルも身につけてきた。前期で学んだ内容を振り返り、後期の学習へとつなげる。

→ これまでの各回の内容を復習しておくこと。

学習到達度の確認試験 前期で学んだ内容の理解度を確認するため、授業時間内に試験を実施する。試験はマークシート形式で行われ、A4用紙2枚までの自筆メモの持ち込みが許可される。テスト理論に基づいて等価性が担保された問題が出題され、前期の学習内容全般から幅広く出題される。特に、記述統計の基本概念、相関と回帰分析の解釈、確率分布の特性、推測統計と帰無仮説検定の考え方など、心理統計の基本的な理解が問われる。

→ 試験の詳細については、別途配布される「試験に関するお知らせ」を参照のこと。

後期の学習に向けて 前期では主にデータの記述と基本的な統計的推論について学んだ。後期ではより複雑な統計モデルや分析手法（分散分析、多変量解析など）について学び、実際の心理学研究に応用する方法を学ぶ。前期で習得した基礎知識をもとに、より高度な統計的思考と分析技術を身につけていく。特に、サンプルから母集団への一般化や、複数の要因が絡む実験計画の分析方法など、心理学研究の実践に直結する内容が中心となる。

→ 夏休み中に前期の内容を復習し、後期の学習に備えること。

13.1.3 キーワード

- 記述統計と推測統計
- 相関と因果
- 確率分布と統計的推論
- 帰無仮説検定
- 統計的思考の実践

13.2 授業情報

■コマの展開方法 前期内容の総括 (約 15 分) の後, 授業時間内試験 (約 60 分) を実施する

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法 (総括)

13.2.1 予習・復習課題

■予習 前期の全授業内容を復習しておくこと。特に以下の点について理解を深めておくことが重要である:

1. 記述統計の基本 (平均, 分散, 標準偏差など) とその意味
2. 相関係数と回帰分析の基本概念と解釈
3. 確率分布 (二項分布, 正規分布など) の特徴と用途
4. 標本分布と中心極限定理の考え方
5. 推測統計の基本概念 (点推定, 区間推定)
6. 帰無仮説検定の論理と手順

また, 授業で扱った R の基本操作とデータ分析の方法についても振り返っておくこと。試験では自筆のメモ (A4 用紙 2 枚まで, 両面可) の持ち込みが許可されるので, 重要な概念や公式をまとめておくといい。

■復習 試験終了後は, 出題された問題について再度考え, 不明点があれば教科書や参考書, 授業ノートで確認しておくこと。また, 後期の学習に向けて, 以下の点を意識しておくといい:

1. 心理学研究における統計的思考の重要性
2. データに基づく客観的な推論の方法
3. 統計ソフトウェア (R) を用いた実践的な分析スキル
4. 心理学の諸問題を統計的に捉える視点

夏休み期間中に前期の内容を総復習し, 特に苦手な分野があれば重点的に学習しておくことを推奨する。後期ではより高度な統計手法を学ぶため, 前期の基礎的な内容をしっかりと理解しておくことが重要である。

14 R を使った t 検定と検定力分析

14.1 授業内容

14.1.1 概要

夏休みを挟んで、前期に学んだ帰無仮説検定の基本的な考え方を復習し、統計ソフトウェア R を用いた実践的な分析方法を学ぶ。特に、心理学研究でよく用いられる t 検定 (対応のある t 検定と対応のない t 検定) に焦点を当て、R による実行方法とその結果の解釈について理解を深める。さらに、シミュレーションを通じて効果量とサンプルサイズが検定結果に与える影響を体験的に学び、適切な研究デザインの重要性について理解する。また、検定力分析の基本的な考え方と実行方法についても学ぶ。

14.1.2 コマ主題細目

帰無仮説検定の復習 前期に学んだ帰無仮説検定の基本的な考え方と手順について復習する。帰無仮説と対立仮説の設定、検定統計量の算出、p 値の解釈、有意水準に基づく判断という一連の流れを確認する。特に、検定結果の解釈と報告方法について理解を深め、統計的有意性と実質的重要性の違いについても再確認する。また、t 分布の特徴と、標準正規分布との関係についても理解する。

→ 川端・莊島 (2014) の P.81–96, 山田・村井 (2004) の P.104–119.

対応のある t 検定 対応のある t 検定 (一要因被験者内計画の 2 水準間の比較) の考え方と実行方法について学ぶ。この検定は同一被験者から得られた 2 つの条件下でのデータ (事前・事後測定など) を比較する場合に用いられる。R での実行方法として、`t.test()` 関数の使い方、特に `paired=TRUE` オプションの意味と使用方法を理解する。また、結果の読み方と解釈、効果量の算出方法についても学ぶ。実際のデータセットを使用した演習を通じて、対応のある t 検定の実践的な実行方法を身につける。

→ 小杉 (2019) の P.126–132, 清水 (2021) の P.124–130.

対応のない t 検定 対応のない t 検定 (独立した 2 群の平均値の比較) の考え方と実行方法について学ぶ。この検定は異なる被験者群から得られたデータを比較する場合に用いられる。R での実行方法として、`t.test()` 関数の基本的な使い方と、分散の等質性を仮定する場合 (`var.equal=TRUE`) と仮定しない場合 (Welch の補正) の違いについて理解する。また、検定結果の読み方と解釈、効果量の算出方法についても学ぶ。実際のデータセットを使用した演習を通じて、対応のない t 検定の実践的な実行方法を身につける。

→ 小杉 (2019) の P.132–138, 清水 (2021) の P.130–136.

効果量とサンプルサイズのシミュレーション R を用いたシミュレーションを通じて、効果量とサンプルサイズが検定結果に与える影響について体験的に学ぶ。異なる効果量 (小・中・大) とサンプルサイズの組み合わせでデータを生成し、t 検定を実行することで、どのような条件で統計的有意性が得られやすいか、または得られにくいかを確認する。これにより、効果量が小さい場合は大きなサンプルサイズが必要であること、サンプルサイズが小さい場合は検出力が低下することなどを理解する。また、シ

ミュレーションの結果から、標本サイズを適切に設定することの重要性についても学ぶ。

→ 効果量については大久保・岡田 (2012) の P.43–46, シミュレーションの基本的な考え方については松村他 (2021) の P.178–186.

検定力分析の基礎 検定力 (第2種の誤りを犯さない確率) の概念と、研究デザインにおける重要性について学ぶ。特に、R の `power.t.test()` 関数を用いた検定力分析の実行方法と解釈について理解する。事前の検定力分析 (研究計画段階での必要サンプルサイズの算出) と事後の検定力分析 (研究結果の評価) の違いと用途について学び、研究デザインの質を高めるための手法として検定力分析を活用する方法を身につける。また、サンプルサイズ、効果量、有意水準、検定力の関係について理解を深める。

→ 検定力分析については川端・荘島 (2014) の P.90–93, 小杉他 (2023) の P.187–224 も参照。

14.1.3 キーワード

- 対応のある t 検定と対応のない t 検定
- R による統計分析の実行
- 効果量と Cohen's d
- サンプルサイズとシミュレーション
- 検定力分析

14.2 授業情報

■コマの展開方法 講義と演習

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; G 推測統計: t 検定

14.2.1 予習・復習課題

■予習 前期に学んだ帰無仮説検定の基本的な考え方と手順について復習しておくこと。特に、帰無仮説と対立仮説の設定、検定統計量の算出、p 値の解釈、有意水準に基づく判断という一連の流れを確認しておくこと。また、R と RStudio の基本的な操作方法についても再確認しておくこと。前期に作成したスクリプトや Rmarkdown ファイルを見直し、基本的なコマンドや関数の使い方を思い出しておくことが望ましい。

■復習 授業で学んだ内容を定着させるために、以下の課題に取り組むこと:

1. 提供されたデータセットを用いて、対応のある t 検定と対応のない t 検定をそれぞれ実行し、結果を Rmarkdown で適切にレポートにまとめる
2. 異なる効果量とサンプルサイズの組み合わせでシミュレーションデータを生成し、t 検定を実行して結果を比較する
3. 与えられた研究シナリオに対して、適切なサンプルサイズを算出するための検定力分析を実行する
4. 検定結果の統計的有意性と効果量の大きさの両方を考慮した解釈を行い、研究における意義について

て考察する

これらの課題を Rmarkdown にまとめ、提出すること。また、検定力分析とサンプルサイズの重要性について理解を深めるために、大久保・岡田 (2012) や川端・荘島 (2014), 小杉他 (2023) の関連する章を読むことが推奨される。

15 実験計画法の基礎と RCT

15.1 授業内容

15.1.1 概要

心理学研究では、変数間の因果関係を検証するために実験的アプローチが重要な役割を果たす。本コマでは、実験計画法の基礎的概念と、特にランダム化比較試験 (Randomized Controlled Trial; RCT) の意義と方法について学ぶ。RCT は因果関係を検証するための強力な方法であり、その理論的背景と実践的手法を理解することは、心理学研究の質を高めるために不可欠である。また、実験における誤差の概念とその分布特性についても学び、統計的推論との関連を理解する。実験計画と統計的分析の橋渡しとなる基本的な考え方を身につけることを目指す。

15.1.2 コマ主題細目

実験計画法の基本的考え方 心理学における実験的アプローチの意義と限界について理解する。実験とは、操作した独立変数が従属変数に与える効果を測定することで因果関係を検証する方法である。相関研究と実験研究の違い、因果関係を主張するための条件 (時間的先行性, 共変関係, 他の説明可能性の排除) について学ぶ。また、実験の内的妥当性と外的妥当性のバランス, 実験における倫理的配慮についても理解する。

→ 三浦 (2017) の P.30–34, 豊田他 (1992) の P.24–66.

ランダム化の意義と方法 ランダム化 (無作為化) は実験研究の核心部分であり、偶然誤差や交絡要因の影響を統制するために不可欠な手法である。ランダム化比較試験 (RCT) では、研究参加者を無作為に実験群と統制群に割り当てることで、群間の初期状態の等質性を確保する。これにより、介入効果以外の要因による影響を最小化し、観察される群間差を介入効果として解釈できるようになる。ランダム化の具体的な方法 (単純ランダム化, ブロックランダム化, 層別ランダム化など) とその利点・欠点について学ぶ。

→ 安井 (2019) を参照, RCT の基本的な考え方については豊田 (2017) の P.5–26.

実験における誤差と正規分布 実験データには必ず誤差 (偶然変動) が含まれる。誤差は様々な要因の集積によって生じるが、中心極限定理により、多くの独立した小さな誤差要因が加算されると、その合計は正規分布に近似する。この性質が実験データの統計的分析の理論的基盤となっている。誤差分布としての正規分布の特性 (平均 0, 対称性, 確率的性質) を理解し、実験計画における誤差の扱い方と統制方法について学ぶ。また、誤差を分散という指標で評価することの意味と、それが実験効果の検出力にどのように影響するかについても理解する。

→ 誤差分布については長沼 (2016) の P.16–40, 山田・村井 (2004) の P.80–85.

実験効果の評価方法 実験効果を評価するためには、介入による変化 (効果) と偶然による変動 (誤差) を区別する必要がある。効果と誤差の分離は、線形モデルの枠組みで理解できる (個々の観測値 = 全体平均 + 介入効果 + 誤差)。効果の大きさは、群間差 (実験群平均 - 統制群平均) で評価するが、

その統計的有意性は誤差の大きさとの相対比較によって判断される。実験効果の評価における統計的検定の役割と、効果量 (標準化された効果の指標) の重要性について理解する。

→ 小杉 (2018) の P.60–67, 効果量については大久保・岡田 (2012) の P.43–46.

実験計画の種類と要素 実験計画を構成する基本要素として、要因 (factor) と水準 (level) の概念を理解する。一要因計画と多要因計画の違い、群間計画 (between-subjects design) と群内計画 (within-subjects design) の特徴と選択基準について学ぶ。また、各種実験計画の表記法 (例:「2×3 の混合計画」) についても理解する。次回以降の講義で詳細に学ぶ群間計画と群内計画の基本的な考え方について概観し、それぞれの利点と欠点を理解する。

→ 豊田 (2017) の P.27–44, 実験計画の種類については小杉 (2018) の P.74–78.

15.1.3 キーワード

- ランダム化比較試験 (RCT)
- 実験群と統制群
- 偶然誤差と正規分布
- 要因と水準
- 群間計画と群内計画

15.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 1. 心理学における実証的研究法 (量的研究及び質的研究); C 実証の手続き

15.2.1 予習・復習課題

■予習 実験計画法の基本的な考え方について、心理学研究法の教科書等で予習しておくこと。特に因果関係と相関関係の違い、実験研究と調査研究の違いについて理解しておくことと講義の理解が深まる。また、前回学んだ t 検定の基本的な考え方を復習しておくこと、実験計画と統計的分析の関連について理解しやすくなる。

■復習 授業で学んだ内容を定着させるために、以下の点について理解を深めておくこと:

1. ランダム化の意義と具体的な方法について説明できるようにする
2. 偶然誤差がなぜ正規分布に従うと考えられるのか、その理論的根拠を理解する
3. 実験効果と誤差を分離する考え方を理解し、線形モデルの枠組みで説明できるようにする
4. 様々な実験計画の種類とその特徴を理解し、具体的な研究例を挙げて説明できるようにする
5. 実験研究の内的妥当性と外的妥当性のバランスについて考察する

また、次回の授業で学ぶ群間計画 (Between-subjects design) の準備として、豊田 (2017) の P.27–44 や小杉 (2018) の P.74–78 を読んでおくことよい。

16 群間計画 (Between Design) と分散分析

16.1 授業内容

16.1.1 概要

前回学んだ実験計画法の基礎知識を踏まえ、本コマでは群間計画 (Between-subjects Design) に焦点を当てる。群間計画とは、異なる実験参加者グループに異なる条件を割り当てる実験デザインであり、心理学研究でよく用いられる手法である。特に、一要因および二要因の群間計画における効果の分解方法と評価方法を学び、分散分析 (ANOVA) の基本的な考え方を理解する。平方和の分解、F 値の算出と解釈、交互作用の概念について学習し、これらが一般線形モデルの枠組みでどのように表現されるかを理解する。データの分析を通じて、研究仮説を統計的に検証する方法を身につける。

16.1.2 コマ主題細目

群間計画の基本構造 群間計画 (Between-subjects Design) は、異なる実験参加者グループに異なる条件を割り当てる実験デザインである。このデザインの特徴、利点と欠点、適用場面について理解する。群間計画では、各参加者は一条件のみを経験するため、練習効果や順序効果といった問題を回避できるが、個人差による誤差が大きくなる傾向がある。一要因計画と多要因計画の違い、具体的な研究例を通じて群間計画の基本構造を理解する。また、群間計画を表記する際の規則 (例:「 2×3 の二要因 Between 計画」) についても学ぶ。

→ 豊田 (2017) の P.27–44, 小杉 (2018) の P.74–78.

平方和の分解と変動の源泉 分散分析の核心は、データの総変動 (平方和) を複数の変動源に分解することにある。一要因群間計画では、総平方和 (SS_T) は要因の効果による平方和 (SS_A) と誤差による平方和 (SS_E) に分解される ($SS_T = SS_A + SS_E$)。二要因群間計画では、さらに第二要因 (SS_B) と交互作用 (SS_{AB}) の平方和が加わる ($SST = SS_A + SS_B + SS_{AB} + SS_E$)。各平方和の計算方法と意味を理解し、効果と誤差をどのように分離するかを学ぶ。また、全体平均、群平均、セル平均の関係と、それらを用いた効果の算出方法についても理解する。

→ 山田・村井 (2004) の P.184–189, 小杉 (2018) の P.81–88.

F 検定とその解釈 平方和を分解した後、効果の統計的有意性を評価するために F 検定を用いる。F 値は、効果の平均平方 ($MS = SS/df$) を誤差の平均平方で割った値 ($F = MS_{\text{効果}}/MS_{\text{誤差}}$) として算出される。自由度 (df) の概念と計算方法、F 分布の特性と F 値の解釈について学ぶ。特に、F 値が大きいほど効果が誤差に比べて相対的に大きいことを意味し、これが統計的に有意であるかどうかを判断するための基準となることを理解する。また、多要因計画における主効果と交互作用それぞれの F 検定と、その結果の解釈方法についても学ぶ。

→ 川端・荘島 (2014) の P.125–130, 山田・村井 (2004) の P.190–196.

交互作用の理解と解釈 二要因以上の実験計画では、要因間の交互作用 (interaction) が重要な意味を持つ。交互作用とは、一方の要因の効果が他方の要因の水準によって異なる現象を指す。交互作用

の種類 (相乗効果, 拮抗効果, クロスオーバー交互作用など) と解釈方法, 交互作用を検出するための統計的手法について学ぶ。また, 交互作用の存在が主効果の解釈にどのような影響を与えるかを理解し, 交互作用がある場合の単純主効果分析など, 適切な追加分析の方法についても学ぶ。交互作用の視覚的理解を助けるグラフの描き方と読み方についても習得する。

→ 山田・村井 (2004) の P.190–193, 小杉 (2019) の P.139–142.

16.1.3 キーワード

- 群間計画 (Between-subjects Design)
- 平方和の分解と分散分析
- F 検定と自由度
- 主効果と交互作用

16.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 2. データを用いた実証的な思考方法; C データの統計的記述

科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; G 推測統計: 代表値と散布度をめぐって (1)

16.2.1 予習・復習課題

■予習 前回学んだ実験計画法の基礎概念 (要因と水準, ランダム化など) を復習しておくこと。また, 線形モデルの基本的な考え方 (切片と傾き, 回帰式) についても再確認しておくことと理解が深まる。さらに, 分散 (variance) の概念と計算方法について復習しておくことも有用である。

■復習 授業で学んだ内容を定着させるために, 以下の点について理解を深めておくこと:

1. 一要因および二要因の群間計画における平方和の分解方法を理解し, 具体例を用いて計算できるようにする
2. 自由度 (df), 平均平方 (MS), F 値の計算方法と意味を理解し, F 検定の結果を適切に解釈できるようにする
3. 交互作用の概念を理解し, 様々なパターンの交互作用 (相乗効果, 拮抗効果, クロスオーバーなど) を図示して説明できるようにする
4. 分散分析を一般線形モデルの枠組みで表現する方法を理解し, 回帰分析との共通点を説明できるようにする
5. 提示された実験デザインを適切に表記し, その構造 (要因, 水準, 群間/群内) を説明できるようにする

次回の授業では, R を用いた群間計画データの分析を行うため, 統計環境 R の基本操作を再確認しておくことが望ましい。特に, データフレームの操作方法や基本的な関数の使い方について復習しておくことと良い。

17 Rによる群間計画の分析

17.1 授業内容

17.1.1 概要

前回学んだ群間計画 (Between-subjects Design) の理論を踏まえ、本コマでは統計環境 R を用いた実践的な分析方法について学ぶ。一要因および二要因の群間計画データを実際に分析し、分散分析表の読み方や結果の解釈方法を習得する。基本的な分析には `aov()` 関数と `lm()` 関数を用い、それらの出力結果から SS(平方和), df(自由度), MS(平均平方), F 値などの情報をどのように読み取るかを学ぶ。また、分析結果を適切に可視化する方法についても実習を通じて理解する。

17.1.2 コマ主題細目

データの準備と基本操作 分散分析を行うためのデータ形式について理解し、適切なデータの読み込み方法と前処理を学ぶ。特に R での分散分析では、カテゴリカル変数を因子型 (factor) として扱うことが重要である。csv ファイルや Excel ファイルからのデータの読み込み方法、データフレームの構造確認、因子型への変換方法、データの要約統計量の算出方法などの基本操作を復習する。また、欠損値の処理方法についても学び、分析前の適切なデータチェックの重要性を理解する。

→ 小杉 (2019) の P.46–70, 松村他 (2021) の P.28–42.

一要因群間計画の分析 R の `aov()` 関数と `lm()` 関数を用いた一要因分散分析の実行方法を学ぶ。公式表記 (formula notation) の書き方 (例: `y ~ group`), 関数の使い方, 出力結果の読み方について理解する。`summary()` 関数を用いた分散分析表 (ANOVA table) の表示方法と、そこに含まれる SS(平方和), df(自由度), MS(平均平方), F 値, p 値の意味を確認する。また, `TukeyHSD()` 関数を用いた多重比較の方法や, 効果量 (η^2 , ω^2 など) の算出方法についても学ぶ。実際のデータセットを用いた演習を通じて, 一要因群間計画の分析手順と結果の解釈を身につける。

→ 小杉 (2019) の P.146–152, 橋本・荘島 (2016) の P.26–55.

二要因群間計画の分析 二要因分散分析の実行方法と結果の解釈について学ぶ。`aov()` 関数における二要因の公式表記 (例: `y ~ factor1 * factor2`) の意味と使い方, 交互作用項の指定方法を理解する。分散分析表から主効果と交互作用の検定結果を読み取る方法, 有意な交互作用がある場合の解釈と追加分析 (単純主効果の検定など) の必要性について学ぶ。また, 二要因分散分析の結果を適切に報告する方法についても理解する。実際のデータセットを用いた演習を通じて, 二要因群間計画の分析手順と結果の解釈を身につける。

→ 小杉 (2019) の P.153–160, 橋本・荘島 (2016) の P.57–90.

分析結果の可視化 分散分析の結果を適切に可視化する方法について学ぶ。R の基本グラフィックス機能や `ggplot2` パッケージを用いた群間の比較を示す棒グラフやボックスプロット, 交互作用を示す交互作用プロットの作成方法を理解する。エラーバーの追加方法, 軸ラベルやタイトルの設定, 凡例の調整など, 読みやすく情報量の多いグラフを作成するためのテクニックを学ぶ。また, 可視化することの重

要性と、異なるグラフが伝える情報の違いについても理解する。実際のデータセットを用いた演習を通じて、分析結果の効果的な可視化方法を身につける。

→ 松村他 (2021) の P.126–138, 小杉 (2019) の P.90–105.

17.1.3 キーワード

- `aov()` 関数と `lm()` 関数
- 分散分析表の読み方 (SS, df, MS, F 値)
- 多重比較と下位検定
- 交互作用の可視化
- `ggplot2` による結果の表現

17.2 授業情報

■コマの展開方法 演習

■標準シラバスにおける位置づけ 科目番号 5 心理学統計法; 2 統計に関する基礎的な知識; C/D/E
エクセル, R, SPSS 入門

17.2.1 予習・復習課題

■予習 前回学んだ群間計画 (Between Design) の基本概念と分散分析の理論を復習しておくこと。特に、平方和の分解, F 値の意味, 主効果と交互作用の概念について理解を確認しておくこと。また, R と RStudio の基本操作, 特にデータの読み込み方法やデータフレームの操作方法について再確認しておくこと。必要に応じて, 第 3 回, 第 4 回の授業で学んだ R 操作についての資料を見直しておくこと。

■復習 授業で学んだ内容を定着させるために, 以下の課題に取り組むこと:

1. 提供されたデータセットを用いて, 一要因群間計画の分散分析を実行し, 結果を適切に解釈する
2. 提供されたデータセットを用いて, 二要因群間計画の分散分析を実行し, 主効果と交互作用の結果を解釈する
3. 分散分析の結果を適切に可視化し, 群間の差や交互作用を示すグラフを作成する
4. `TukeyHSD()` 関数を用いて多重比較を行い, どの群間に有意差があるかを特定する
5. 上記の分析結果を Rmarkdown を用いてレポート形式にまとめる

これらの課題を Rmarkdown ファイルにまとめ, 期日までに提出すること。時間内に終わらなかった場合は, 次回授業までに完成させておくこと。分析結果の適切な解釈と可視化を重視し, 統計的に有意な結果だけでなく, 効果量や信頼区間についても言及するよう心がけること。不明点があれば, 小杉 (2019) や橋本・荘島 (2016) を参照するか, オフィスアワーを利用して質問すること。

18 群内計画 (Within Design) の理論

18.1 授業内容

18.1.1 概要

これまで学んだ群間計画に続き、本コマでは群内計画 (Within-subjects Design) について学ぶ。群内計画は、同一の実験参加者が複数の条件を経験する実験デザインであり、事前-事後測定や反復測定などがこれに該当する。群内計画の最大の特徴は、個人差を個別に識別し、誤差から分離できることである。このことにより分析の精度が高まり、効果の検出力が向上する。本コマでは、群内計画の基本的な考え方、データの構造、効果と誤差の分離方法、個人差の取り扱い方について理解し、さらに階層モデルとしての統計的特性についても学ぶ。

18.1.2 コマ主題細目

群内計画の基本構造 群内計画 (Within-subjects Design) は、同一の実験参加者が複数の条件を経験する実験デザインである。このデザインの特徴、利点と欠点、適用場面について理解する。群内計画の代表的な例として、事前-事後測定、反復測定、対応のある比較などがある。個人差を誤差から分離できるという最大の利点がある一方で、練習効果、疲労効果、順序効果といった問題も生じうることを理解する。また、実験のコストや実施上の制約、倫理的配慮についても考察する。群内計画を表記する際の規則 (例:「3 水準の一要因 Within 計画」) についても学ぶ。

→ 豊田 (2017) の P.45–62, 小杉 (2018) の P.74–78.

個人差の分離と精度の向上 群内計画の最大の特徴は、個人差を誤差から分離できることである。従来の群間計画では、個人差は誤差に含まれていたが、群内計画では個人要因として明示的に扱うことができる。これにより比較すべき誤差 (残差) が小さくなり、効果の検出力が向上する。個人差の分散と残差の分散を分離する方法、それによって F 値がどのように変化するかを理解する。また、効果量の算出方法と解釈についても学ぶ。視覚的な例を通じて、個人差を分離することの意義を直感的に理解する。

→ 山田・村井 (2004) の P.178–183, 小杉 (2019) の P.146–148.

群内計画のモデルと平方和の分解 群内計画のデータ構造を数学的に表現し、平方和の分解方法を理解する。総平方和 (SS_T) は、要因の効果による平方和 (SS_A)、個人差による平方和 (SS_S)、および残差による平方和 (SS_E) に分解される ($SS_T = SS_A + SS_S + SS_E$)。各平方和の計算方法と意味を理解し、効果と誤差、個人差をどのように分離するかを学ぶ。また、適切な比較対象となる平方和 (誤差項) の選択方法と、F 値の算出方法についても理解する。群内計画の分散分析表の読み方と、結果の解釈方法についても学ぶ。

→ 山田・村井 (2004) の P.179, 図 7.5.1 の平方和の分解図式, 小杉 (2018) の P.81–88.

群内計画と階層モデル 群内計画は、統計的には階層モデル (multilevel model) あるいは混合モデル (mixed model) として捉えることができる。個人差という別の分布が入った統計モデルとして理解す

ることで、より複雑な実験デザインにも対応できる柔軟性を持つ。固定効果 (fixed effect) と変量効果 (random effect) の違い、個人要因を変量効果として扱う意義、変量効果モデルの特性について学ぶ。また、個人差が正規分布に従うという仮定の妥当性と、その検証方法についても理解する。階層モデルの考え方は、より複雑な統計モデル (一般化線形混合モデルなど) への橋渡しとなることについても触れる。

→ 小杉 (2018) の P.136–138, 階層モデルについては豊田 (2017) の P.121–145.

18.1.3 キーワード

- 群内計画 (Within-subjects Design)
- 個人差の分離と平方和の分解
- 反復測定と対応のある比較
- 階層モデルと混合モデル
- 固定効果と変量効果

18.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 2. データを用いた実証的な思考方法; C データの統計的記述

科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; G 推測統計: 代表値と散布度をめぐって (1)

18.2.1 予習・復習課題

■予習 前回までに学んだ群間計画 (Between Design) の特徴と分散分析の基本的な考え方を復習しておくこと。特に、一要因の分散分析における平方和の分解、F 値の意味について理解を確認しておくこと。また、実験計画の用語 (要因、水準など) についても再確認しておくこと。心理学の論文から群内計画を用いた研究例を探して読んでみることも有用である。

■復習 授業で学んだ内容を定着させるために、以下の点について理解を深めておくこと:

1. 群間計画と群内計画の違い、それぞれの利点と欠点を説明できるようにする
2. 群内計画における平方和の分解方法を理解し、個人差が誤差から分離される過程を説明できるようにする
3. 群内計画における適切な F 値の算出方法を理解し、分散分析表を正しく解釈できるようにする
4. 個人差を変量効果として扱う階層モデル・混合モデルの考え方を理解し、その特徴を説明できるようにする
5. 実際の研究例をもとに、群内計画が適切な状況とそうでない状況を判断できるようにする

群内計画における個人差を算出するための手続きを、具体的な数値例を用いて自分で確認してみるとよい。公開されている実験データや、教科書に載っている例題を使って練習するとよい。また、混合モデル (階層線形モデル、マルチレベルモデルとも呼ばれる) については、さまざまな呼称があるため、それぞれのキー

ワードで参考文献を探してみるとよい。

19 R による Within デザインの分散分析

19.1 授業内容

19.1.1 概要

前回学んだ群内計画 (Within-subjects Design) の理論的理解をもとに、本コマでは R を用いた群内計画の実践的な分析方法について学ぶ。特に、`anovakun` という関数群を活用し、群内計画データの構造化、モデルの構築、分散分析の実施、結果の解釈までの一連のプロセスを習得する。群内計画特有の個人差の分離、効果の検出、交互作用の解釈について R の出力結果から理解を深める。また、分析結果の適切な可視化と報告方法についても学び、心理統計学の実践的スキルを向上させる。

19.1.2 コマ主題細目

群内計画データの構造と R での表現 群内計画のデータ構造を R で表現する方法について学ぶ。`long` 形式(各行が 1 観測値)と `wide` 形式(各行が 1 参加者のすべての条件データ)の違いと変換方法、群内計画においてはどちらの形式が分析に適しているかを理解する。また、R における因子型変数 (factor) の設定と水準の指定方法、参加者 ID 変数の扱い方についても学ぶ。`anovakun` パッケージの基本的な機能とインストール方法、読み込み方法も確認する。

→ 豊田 (2017) の P.215–220, `anovakun` の解説は **aoki2021** の P.128–132.

`anovakun` による一要因群内計画の分析 `anovakun` パッケージを用いた一要因群内計画の分析方法を習得する。ANOVA 関数と群内計画を指定するための構文(参加者 ID 変数の指定方法など)、モデル式の記述法を理解する。得られた分散分析表の読み方と解釈、特に個人差(参加者間変動)が明示的に分離されていることを確認する。効果量(η^2 , ω^2 など)の算出方法と、分析結果の適切な報告形式についても習得する。

→ **aoki2021** の P.133–142, 小杉 (2019) の P.152–157.

`anovakun` による多要因群内計画の分析 二要因以上の群内計画の分析方法について学ぶ。複数の実験要因がすべて群内要因である場合のモデル式の記述法、交互作用の指定方法、結果の解釈方法を理解する。特に、交互作用が有意であった場合の単純主効果の検定方法について `anovakun` での実施方法を習得する。また、多要因群内計画における効果の分解と誤差項の選択、適切な F 値の算出方法についても学ぶ。実際の心理学実験データを用いた演習を通じて、多要因群内計画の分析スキルを向上させる。

→ **aoki2021** の P.143–156, 小杉 (2018) の P.125–132.

群内計画の結果の可視化と報告 群内計画の分析結果を適切に可視化し報告する方法について学ぶ。`ggplot2` パッケージを用いた群内計画特有のエラーバーの設定(個人内変動を反映した信頼区間など)、交互作用の視覚的表現方法、`ggplot2` での作図コードと解釈について理解する。また、APA 形式に準拠した統計結果の報告方法、効果量の適切な提示方法、群内計画特有の注意点(個人差の分離に言及するなど)についても学ぶ。

→ aoki2021 の P.157–165, 小杉 (2019) の P.160–165.

19.1.3 キーワード

- anovakun パッケージによる群内計画分析
- long 形式と wide 形式のデータ構造
- 一要因・多要因群内計画の実装
- 個人差の分離と効果量の算出
- ggplot2 による群内計画結果の可視化

19.2 授業情報

■コマの展開方法 講義と R 実習

■標準シラバスにおける位置づけ 科目番号 4; 心理学研究法; 2. データを用いた実証的な思考方法; C データの統計的記述

科目番号 5; 心理学統計法; 1. 心理学で用いられる統計手法; G 推測統計: 代表値と散布度をめぐって (1)

19.2.1 予習・復習課題

■予習 前回学習した群内計画 (Within Design) の理論的内容を復習しておくこと。特に、群内計画における個人差の分離の意義、平方和の分解、適切な F 値の算出方法について理解を確認しておくこと。また、R の基本操作、データフレームの扱い方、基本的な関数の使用方法についても再確認しておくこと。

■復習 授業で学んだ内容を定着させるために、以下の点について理解を深めておくこと:

1. 群内計画データの適切な構造化方法を理解し、wide 形式と long 形式の相互変換ができるようにする
2. anovakun パッケージを用いた一要因群内計画の分析を独力で実施し、結果を正しく解釈できるようにする
3. 多要因群内計画の分析方法を理解し、交互作用の解釈と単純主効果の検定ができるようにする
4. 群内計画の分析結果を適切に可視化し、APA 形式に従った報告ができるようにする
5. 実際の研究データや公開データを用いて、群内計画の分析を練習し、スキルを向上させる

授業で扱った分析スクリプトを、異なるデータセットに適用してみることで理解を深めるとよい。公開されている心理実験データや教科書に掲載されているデータを用いて実習してみるとよい。また、各自の研究計画において群内計画が適用可能かどうかを検討し、必要に応じて分析計画を立ててみるとよい。なお、anovakun の開発者の Web サイトや GitHub ページで提供されている追加の解説やチュートリアルも参考になる。

20 ベイズ推定による回帰分析

20.1 授業内容

20.1.1 概要

前回学んだベイズ統計学の基本原理をもとに、本コマではベイズ推定を回帰分析に適用する方法について学ぶ。回帰モデルの各パラメータ(切片と傾き)を確率変数として扱い、そのベイズ推定法と解釈について理解する。また、事後分布から得られる推定値の種類(EAP, MAP, MED など)と確信区間の意味、さらに複数のモデルを比較するための周辺尤度とベイズファクターの概念を学ぶ。これにより、頻度主義的な回帰分析とベイズ的アプローチの共通点と相違点を理解し、それぞれの利点を活かした統計的推論ができるようになることを目指す。

20.1.2 コマ主題細目

回帰モデルのベイズ表現 単回帰モデル $Y_i = \beta_0 + \beta_1 X_i + \varepsilon_i$ を確率モデルとして表現し、ベイズ推定の枠組みで捉え直す。誤差項 ε_i が正規分布 $N(0, \sigma^2)$ に従うと仮定する従来の回帰モデルでは、パラメータ $\beta_0, \beta_1, \sigma^2$ は未知の固定値として扱われるが、ベイズアプローチではこれらのパラメータを確率変数として扱う。具体的には、各パラメータに事前分布を設定し(例えば β_0, β_1 に正規分布, σ^2 に逆ガンマ分布など)、データが与えられた下での事後分布 $p(\beta_0, \beta_1, \sigma^2 | data)$ を求める。ベイズの定理により、この事後分布は事前分布と尤度関数の積に比例することを理解する。ベイズ回帰モデルは、パラメータの不確実性を直接モデル化できるというメリットを持つことを学ぶ。

事後分布からの推定値と確信区間 ベイズ推定では、事後分布から様々な推定値を導出できることを学ぶ。代表的な推定値として、事後平均(EAP: Expected A Posteriori), 事後最頻値(MAP: Maximum A Posteriori), 事後中央値(MED: Median)がある。EAP は事後分布の期待値であり、二乗損失関数の下で最適な推定値となる。MAP は事後分布のモードであり、最尤推定値と密接に関連する(無情報事前分布の場合、MAP と最尤推定値は一致する)。MED は事後分布の中央値であり、絶対損失関数の下で最適となる。また、事後分布からは 95% 確信区間(Credible Interval)を導出できるが、これは頻度主義的な 95% 信頼区間とは解釈が異なることを理解する。確信区間は「パラメータが区間内に入る確率が 95% である」と直接解釈できる点が特徴である。さらに、最高密度区間(HDI: Highest Density Interval)という、事後分布の密度が最も高い領域に基づく区間についても学ぶ。

事前分布の選択と感度分析 回帰モデルのベイズ推定における事前分布の役割と選択方法について学ぶ。無情報事前分布から強い情報を持つ事前分布まで、様々な選択肢とその影響を理解する。事前分布の選択は、サンプルサイズが小さい場合に特に重要であり、事前情報の強さによって事後分布が大きく変わる可能性がある。このため、異なる事前分布のもとでの推定結果の安定性を確認する「感度分析」の重要性を学ぶ。無情報事前分布を用いた場合でも、ベイズ推定は最尤推定とは異なる結果を与えることがあり(特に複雑なモデルや少数のデータの場合)、その違いの原因と解釈について理解する。また、事前分布の選択と報告の透明性の重要性についても学ぶ。

モデル比較と周辺尤度 ベイズ統計学におけるモデル比較の方法として、周辺尤度(marginal likelihood)とベイズファクター(Bayes Factor)の概念を学ぶ。周辺尤度 $p(data|Model)$ は、あるモデルの下

でデータが観測される確率であり、ベイズの定理における分母(正規化定数)に相当する。これは $p(\text{data}|\text{Model}) = \int p(\text{data}|\theta, \text{Model})p(\theta|\text{Model})d\theta$ と表される。つまり、パラメータ空間全体にわたる尤度と事前分布の積の積分として計算される。二つのモデル M_1 と M_2 を比較する場合、ベイズファクター $BF_{12} = \frac{p(\text{data}|M_1)}{p(\text{data}|M_2)}$ がモデル M_1 の M_2 に対する相対的な証拠の強さを表す。ベイズファクターは、帰無仮説検定における p 値とは異なり、帰無仮説を積極的に支持する証拠を提供できる点が特徴である。複雑なモデルでは周辺尤度の計算が困難であるため、ネストしたモデルの比較にはサヴェージ・ディッキー比(Savage-Dickey ratio)などの近似法が用いられることを簡単に紹介する。

→ Lee・Wagenmakers (2013 井関訳 2017) の P.88–103, 豊田 (2015) の P.56–60.

20.1.3 キーワード

- ベイズ回帰モデル
- 事後分布からの推定値(EAP, MAP, MED)
- 確信区間と最高密度区間(HDI)
- 周辺尤度とベイズファクター
- 事前分布の選択と感度分析

20.2 授業情報

■コマの展開方法 講義

■標準シラバスにおける位置づけ 科目番号 5;心理学統計法; (1) 心理学で用いられる統計手法; J より高度な記述や推定を目指して

20.2.1 予習・復習課題

■予習 前回学んだベイズ統計学の基本概念(ベイズの定理, 事前分布, 事後分布など)を復習しておくこと。また, 第 5, ??, 7 講で学んだ相関分析と回帰分析の基本についても再確認しておくこと。特に, 回帰モデルにおけるパラメータの意味, 最小二乗法と最尤法による推定方法の違いについて理解を深めておくことで, ベイズ推定との比較がしやすくなる。

■復習 授業で学んだ内容を定着させるために, 以下の点について理解を深めておくこと:

1. 回帰モデルをベイズ的な枠組みで表現し, 事前分布と事後分布の関係を説明できるようにする
2. 事後分布から得られる様々な推定値(EAP, MAP, MED など)の違いと特徴を理解する
3. 頻度主義的な信頼区間とベイズ的な確信区間の解釈の違いを明確に説明できるようにする
4. 事前分布の選択が事後分布にどのように影響するかを理解し, 適切な事前分布を選択する基準を考える
5. 周辺尤度とベイズファクターの計算方法と解釈を理解し, モデル比較に活用できるようにする

また, 簡単な例(例: 少数のデータを用いた単回帰モデル)を自分で考え, 異なる事前分布を設定した場合

の事後分布の変化を考察してみるとよい。次回以降の授業では JASP などのソフトウェアを用いたベイズ推定の実践を行うため、今回の概念的理解をしっかりと固めておくことが重要である。

21 JASP によるベイズ分析入門

21.1 授業内容

21.1.1 概要

これまで理論的に学んできたベイズ統計学を実践するために、オープンソースの統計ソフトウェア JASP (Jeffreys's Amazing Statistics Program) を導入する。JASP は直感的な GUI を備え、従来の頻度主義的手法とベイズ統計学的手法を並列的に提供する特徴を持つ。本コマでは、JASP のインストールと基本的な操作方法を学び、サンプルデータを用いた記述統計、可視化、相関分析、回帰分析を実践する。特に、従来の頻度主義的分析とベイズ的分析の結果を比較しながら、両者の出力の違いと解釈方法について理解を深める。これにより、統計ソフトウェアを用いたベイズ分析の基礎を身につけ、次回以降のより高度な分析への準備とする。

21.1.2 コマ主題細目

JASP の導入とインターフェイス JASP はオープンソースの統計ソフトウェアで、心理学研究のための統計分析機能と直感的なユーザーインターフェイスを提供する。まず、JASP のウェブサイト (<https://jasp-stats.org/>) からソフトウェアをダウンロードし、各自の PC にインストールする方法を学ぶ。続いて、JASP の基本的なインターフェイスと機能(メニュー構成、データビュー、分析ビュー、結果ビュー)について理解する。JASP の特徴として、オープンサイエンスをサポートする機能(分析結果の再現性、OSF との連携)や、従来の頻度主義的分析とベイズ的分析を同じメニューから実行できる利便性について学ぶ。また、サンプルデータの読み込み方、外部データ(CSV, SPSS など)のインポート方法、変数の種類の設定方法についても実習を通じて習得する。

→ 清水・山本 (2022) の P.1–15.

記述統計と可視化 JASP を用いた基本的なデータ分析として、記述統計と可視化の方法を学ぶ。「Descriptives」メニューを使って、各変数の基本統計量(平均、中央値、標準偏差、四分位数など)を計算し、結果の読み方を理解する。また、ヒストグラム、箱ひげ図などのグラフを作成し、データの分布特性を視覚的に把握する方法を学ぶ。これらの基本操作を通じて、JASP のワークフローに慣れ、次のステップであるより高度な分析への準備を整える。

→ 清水・山本 (2022) の P.1–15.

相関分析：頻度主義的手法とベイズ的手法 JASP を用いた相関分析を、伝統的手法とベイズ的手法の両方で実施し、結果を比較する。頻度主義的アプローチでは、相関係数、 p 値、信頼区間が出力される一方、ベイズ的アプローチでは、ベイズファクター(BF_{10})、事後分布、確信区間が出力される違いを理解する。特に、帰無仮説(相関がない)と対立仮説(相関がある)の間のベイズファクターが、証拠の強さをどのように表現するかを学ぶ。

回帰分析のベイズ推定 JASP を用いた回帰分析を、頻度主義的アプローチとベイズ的アプローチの両方で実施する。「回帰」メニューの「線形回帰」を伝統的手法とベイズ的手法、それぞれ使用して分析を行い、結果の違いを比較する。ベイズ回帰分析では、回帰係数の事後分布、確信区間、モデル比較のた

めのベイズファクターが出力される。特に、従来の p 値に基づく変数選択と、ベイズファクターに基づくモデル比較の違いを理解する。さらに、次回の内容につながる準備として、 t 検定や分散分析のベイズ版についても簡単に触れる。

21.1.3 キーワード

- JASP (Jeffreys's Amazing Statistics Program)
- オープンソース統計ソフトウェア
- ベイズファクター (BF_{10})
- 事後分布の可視化

21.2 授業情報

■コマの展開方法 実習

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; (1) 心理学で用いられる統計手法; J より高度な記述や推定を目指して

21.2.1 予習・復習課題

■予習 授業前に JASP をダウンロードし、各自の PC にインストールしておくこと (<https://jasp-stats.org/> からダウンロード可能)。インストール後、ソフトウェアが正常に起動するか確認しておくこと。また、これまでの授業で学んだベイズ統計学の基本概念(事前分布, 事後分布, ベイズファクターなど)を復習しておくこと。さらに、第 5, ?? 講で学んだ相関分析と回帰分析の基本についても再確認しておくこと。

■復習 授業で学んだ内容を定着させるために、以下の点について実践を通じて理解を深めておくこと:

1. JASP の基本的な操作(データの読み込み, 変数タイプの設定, 分析の実行, 結果の保存)を繰り返し練習する
2. 自分で入手したデータや, JASP に内蔵されているサンプルデータを用いて, 記述統計と可視化を実行してみる
3. 相関分析を頻度主義的アプローチとベイズ的アプローチの両方で実行し, 結果の違いを比較検討する
4. 回帰分析についても同様に両アプローチで実行し, 係数の推定値, 信頼/確信区間, p 値/ベイズファクターの違いをまとめる

特に, JASP の操作に慣れることが重要であるため, 様々なデータセットを用いて分析を繰り返し実行してみるとよい。また, JASP の公式ウェブサイトには多数のチュートリアルやガイドが掲載されているので, それらを参考にしながら学習を進めることも有効である。次回の授業では検定や分散分析のベイズ版を扱うため, 基本操作を確実に身につけておくことが重要である。

22 ベイズ的判断とベイズファクター

22.1 授業内容

22.1.1 概要

前回 JASP の基本的な使用方法を学んだ上で、本コマではベイズ統計学における判断方法について理解を深める。まず、事後分布に基づく確信区間(Credible Interval)と実質的等価性の領域(ROPE: Region of Practical Equivalence)を用いた判断方法について学ぶ。次に、モデル比較の手法としてのベイズファクターの概念と解釈、そして周辺尤度の計算方法としてのサヴェージ・ディッキー法について理解する。特に、従来の帰無仮説検定では帰無仮説を棄却することしかできなかったのに対し、ベイズ法では帰無仮説を積極的に支持することも可能である点を強調する。最後に、情報仮説を含むより柔軟なモデル比較や、実践的なデータ分析におけるベイズ的アプローチの利点について学ぶ。これらを JASP を用いた実例を通じて理解することで、より洗練された統計的判断力を身につける。

22.1.2 コマ主題細目

確信区間と ROPE に基づく判断 ベイズ推定では、パラメータの事後分布から確信区間(Credible Interval)を算出できることを学ぶ。例えば 95% 確信区間は、「パラメータが区間内に 95% の確率で存在する」と直接的に解釈できる。これは頻度主義的な信頼区間とは解釈が異なり、より直感的である。また、実質的等価性の領域(ROPE: Region of Practical Equivalence)の概念を導入し、パラメータがゼロや特定の値と「実質的に等価」かどうかを判断する方法を学ぶ。具体的には、95% 確信区間が ROPE に完全に含まれる場合は「実質的に等価」、完全に外れる場合は「実質的に異なる」、部分的に重なる場合は「判断保留」と判断できる。これにより、単なる「有意/非有意」の二分法ではなく、パラメータの大きさと不確実性を考慮した判断が可能になることを理解する。JASP を用いて、実際のデータセットに対するこれらの判断方法を実演する。

ベイズファクターの概念と解釈 ベイズ統計学における仮説検証の方法として、ベイズファクター(Bayes Factor, BF)の概念と解釈について学ぶ。ベイズファクターは、二つのモデルの相対的な証拠の強さを表す比率であり、 $BF_{10} = \frac{p(data|H_1)}{p(data|H_0)}$ と定義される。ここで $p(data|H)$ は、モデル H のもとでデータが観測される確率(周辺尤度)である。 $BF_{10} > 1$ のとき対立仮説 H_1 が、 $BF_{10} < 1$ のとき帰無仮説 H_0 が支持される。 BF_{10} の値の大きさは証拠の強さを表し、例えば $BF_{10} = 10$ は「対立仮説のための証拠が帰無仮説のための証拠の 10 倍ある」ことを意味する。ベイズファクターの解釈基準(例: 1-3 は「わずかな証拠」、3-10 は「中程度の証拠」、10 以上は「強力な証拠」)について学び、具体的な研究例での解釈方法を理解する。JASP では、分析結果にベイズファクターが自動的に含まれることを確認し、その読み方を実践的に学ぶ。

→ Lee・Wagenmakers (2013 井関訳 2017) の P.88-103.

周辺尤度とサヴェージ・ディッキー法 ベイズファクターを計算するには、各モデルの周辺尤度 $p(data|H)$ を求める必要がある。周辺尤度は、パラメータ空間全体にわたる尤度と事前分布の積の積分 $p(data|H) = \int p(data|\theta, H)p(\theta|H)d\theta$ として定義される。しかし、複雑なモデルではこの積分の計算が困難であるため、近似計算法が必要となる。特に、ネストした(入れ子になった)モデル比

較の場合、サヴェージ・ディッキー法(Savage-Dickey Density Ratio Method)が有用である。これは、対立仮説の事前分布と事後分布のパラメータ特定値(例えばゼロ)における比率からベイズファクターを近似する方法である。JASP では、t 検定や分散分析などのベイズ版で、このサヴェージ・ディッキー法が内部的に使用されていることを理解する。実例として、JASP を用いた単純な t 検定と分散分析のベイズ版を実行し、出力されるベイズファクターの解釈方法を学ぶ。

→ Kruschke (2014 前田・小杉監訳 2017) の P.342–364, Lee・Wagenmakers (2013 井関訳 2017) の P.104–120.

情報仮説と実践的応用 ベイズ統計学の大きな利点として、帰無仮説対対立仮説という二分法だけでなく、より洗練された「情報仮説(informed hypothesis)」の検証が可能である点を学ぶ。例えば、パラメータが特定の範囲内にあるという仮説や、複数のパラメータ間の大小関係に関する仮説など、研究者の事前知識や理論に基づいた具体的な仮説を検証できる。これにより、単なる「効果があるかないか」ではなく、「理論と一致する効果があるか」をより直接的に検証できる。JASP ではこのような情報仮説の検証も可能であり、その設定方法と解釈を学ぶ。また、ベイズ統計学の実践的応用として、サンプルサイズが小さい場合の推測、複数の実験結果の統合、事前知識の体系的な組み込みなど、従来の頻度主義的アプローチでは難しかった分析が可能になることを理解する。具体的な研究例を通じて、ベイズ的アプローチが心理学研究にもたらす新たな可能性に言及する。

22.1.3 キーワード

- 確信区間(Credible Interval)
- 実質的等価性の領域(ROPE)
- ベイズファクター(Bayes Factor)
- サヴェージ・ディッキー法
- 情報仮説(Informed Hypothesis)

22.2 授業情報

■コマの展開方法 講義と JASP を用いた実習

■標準シラバスにおける位置づけ 科目番号 5;心理学統計法; (1) 心理学で用いられる統計手法; J より高度な記述や推定を目指して

22.2.1 予習・復習課題

■予習 前回学んだ JASP の基本操作を復習し、スムーズに使えるようにしておくこと。また、ベイズ統計学の基本概念(事前分布, 事後分布, ベイズの定理など)についても再確認しておくこと。特に、従来の帰無仮説検定のロジック(帰無仮説が正しいという仮定のもとで、観測されたデータが得られる確率を評価する)と、ベイズ的アプローチのロジック(データが与えられたときの各モデルの相対的な確からしさを評価する)の違いについて理解を深めておくことが重要である。

■復習 授業で学んだ内容を定着させるために、以下の点について理解を深めておくこと:

1. 確信区間と ROPE を用いた判断方法の手順と解釈を自分の言葉で説明できるようにする
2. ベイズファクターの計算方法, 解釈基準, BF_{10} と BF_{01} の関係を理解する
3. サヴェージ・ディッキー法の基本原理と, なぜこの方法が周辺尤度の計算を簡略化できるかを説明できるようにする
4. 情報仮説の設定方法と, それが従来の帰無仮説/対立仮説の枠組みとどのように異なるかを理解する
5. JASP を用いて実際のデータセットに対するベイズ的判断を実施し, 結果を適切に解釈できるようにする

特に, JASP 内蔵のサンプルデータセットや, 自分で入手した心理学的データを用いて, t 検定, 分散分析, 回帰分析などのベイズ版を実行し, ベイズファクターや確信区間の解釈を練習するとよい。また, 同じデータに対して頻度主義的アプローチとベイズ的アプローチの両方を適用し, 結果の違いとその理由について考察することで, 両者の特徴をより深く理解することができるだろう。ベイズ統計学は心理学研究における新たな標準となりつつあるため, 実践的なスキルを身につけることが重要である。

23 ベイジアンモデリングの展開

23.1 授業内容

23.1.1 概要

これまでの授業でベイズ統計学の基本原理と実践方法を学んできたが、本コマではベイジアンモデリングの多様な発展的応用と可能性について理解を深める。従来の統計学では扱いにくかった様々な複雑なデータ構造や研究課題に対して、ベイズ統計学と確率的プログラミングの組み合わせがどのように新たな解決策を提供するかを学ぶ。具体的には、分散の検定、打ち切りデータの取り扱い、捕獲再捕獲問題のような計数データのモデリング、複数の分布が混合した混合分布モデルなどの例を通じて、ベイジアンモデリングの柔軟性と表現力の豊かさを理解する。本コマは心理統計学の基礎コースの総括として、統計モデリングの考え方がいかに心理学研究の問いに対応できるかを展望する。

23.1.2 コマ主題細目

ベイジアンモデリングの基礎と確率的プログラミング ベイジアンモデリングとは、データを生成する確率過程を明示的にモデル化し、ベイズの定理を用いてそのモデルのパラメータを推定する方法である。このアプローチの特徴は、複雑な確率モデルを柔軟に構築できる点にある。確率的プログラミング言語(Stan, JAGS, PyMC など)の登場により、これらの複雑なモデルの推定が実用的になった。確率的プログラミングでは、モデルの構造を直接的に記述し、MCMC などの手法を用いて事後分布を数値的に近似することができる。これにより、頻度主義的アプローチでは扱いにくかった複雑なモデルの推定が可能になったことを理解する。特に、階層モデル、混合モデル、非線形モデルなど、心理学研究で重要な様々なモデルが実装できるようになった歴史的背景と意義について学ぶ。

→ Lee・Wagenmakers (2013 井関訳 2017) の P.3–13

分散のモデリングと検定 従来の統計学では、平均値に関する推測が中心であり、分散の検定は比較的扱いが難しかった。ベイジアンモデリングでは、分散自体を確率変数として直接モデル化できるため、分散に関する推測がより直接的に行える。例えば、二群の分散を比較する問題や、条件によって反応のばらつきが変化するかどうかを検討する問題などが容易に扱える。心理学研究では、平均だけでなく分散(個人差や反応のばらつき)も重要な情報を含むことが多いため、これらの手法の有用性は高い。

→ Lee・Wagenmakers (2013 井関訳 2017) の P.48–51.

打ち切りデータと特殊な分布のモデリング 心理学研究では、打ち切りデータ(censored data)や切断データ(truncated data)など、標準的な統計手法では扱いにくいデータ構造がしばしば現れる。例えば、反応時間データでタイムアウトがある場合や、質問紙で天井効果・床効果が生じる場合などである。ベイジアンモデリングでは、これらの特殊なデータ生成過程を明示的にモデル化できる。具体的には、打ち切りや切断を含む確率分布を定義し、観測されたデータがその一部であるという条件の下で推測を行う方法を学ぶ。また、超幾何分布やディリクレ分布など、特殊な確率分布を用いたモデリングについても理解する。例として、捕獲再捕獲問題(capture-recapture problem)をベイジアンモデリ

ングで解く方法を紹介し、母集団サイズの推定など、従来は扱いにくかった問題への応用を学ぶ。

→ Lee・Wagenmakers (2013 井関訳 2017) の P.61–66.

混合分布モデルと潜在クラスモデル 心理学のデータは、しばしば複数の異なるプロセスや集団が混在した結果である。例えば、実験参加者の中に課題を理解している人とそうでない人が混在している場合や、反応時間が複数の認知プロセスの組み合わせを反映している場合などである。混合分布モデル (mixture model) は、このような状況を「複数の分布の重み付き和」としてモデル化する手法である。ベジアンアプローチは、混合分布モデルの推定に特に適している。具体例として、反応時間データの混合モデル(例: ex-Gaussian 分布)や、潜在クラスモデル(latent class model)によるグループ分けなどを取り上げる。これらのモデルは、単一の分布では捉えきれない複雑なデータ構造を理解するのに役立つ。また、潜在変数の導入により、直接観測できない心理プロセスを推測する方法についても学ぶ。

23.1.3 キーワード

- ベイジアンモデリング
- 確率的プログラミング
- 分散の検定とモデリング
- 打ち切りデータと特殊分布
- 混合分布モデルと潜在クラス

23.2 授業情報

■コマの展開方法 講義と事例紹介

■標準シラバスにおける位置づけ 科目番号 5; 心理学統計法; (1) 心理学で用いられる統計手法; J より高度な記述や推定を目指して

23.2.1 予習・復習課題

■予習 これまでの授業で学んだベイズ統計学の基本概念と実践方法を復習しておくこと。特に、ベイズの定理、事前分布と事後分布の関係、ベイズファクターなどの基本的な概念について再確認しておくこと。また、前回までの授業で扱った分析例について振り返り、ベイズ的アプローチの特徴と利点について理解を深めておくことが重要である。

■復習 授業で学んだ内容を定着させ、発展的な理解を深めるために、以下の点について考察し、整理しておくこと:

1. ベイジアンモデリングがどのように複雑な研究課題やデータ構造に対応できるか、具体例を挙げて説明できるようにする
2. 自分の研究分野や関心のある心理学的問題に対して、ベイジアンモデリングをどのように適用できるかを考察する
3. 分散のモデリング、打ち切りデータの扱い、混合分布モデルなど、授業で紹介された具体的な事例に

ついて、その基本原理と心理学研究への応用可能性を理解する

4. 確率的プログラミングの基本的な考え方と、それが従来の統計ソフトウェアとどのように異なるかを説明できるようにする
5. 本授業で学んだ統計的概念とモデリングの考え方が、心理学の理論構築と実証にどのように貢献できるかを総合的に考察する

また、更なる学習のために、授業で紹介された文献や参考図書、オンライン資料などに目を通し、興味を持った特定のトピックについて理解を深めるとよい。ベイジアンモデリングは急速に発展している分野であり、今後の心理学研究において重要な役割を果たすことが期待されるため、基本的な理解を固めておくことが重要である。

24 授業内試験

24.1 授業内容

24.1.1 概要

本コマでは、これまでに学んできた心理統計学の内容についての理解度を確認するための授業内試験を実施する。前期から後期にわたる授業で学んだ基本的な知識、重要な統計概念、データ分析の方法論、そして心理統計学の考え方の筋道について総合的に問う試験を通じて、一年間の学習の成果を測定する。試験の内容は、記述統計から推測統計、回帰分析、分散分析、そして近年重要性を増しているベイズ統計学までを含む幅広いものとなる。試験結果は成績評価の重要な要素となるため、これまでの学習内容を総復習し、心理統計学の基礎を確実に身につけておくことが望ましい。

24.1.2 コマ主題細目

試験の概要と出題範囲 授業内試験はマークシート形式で実施され、問題数は 60～70 問程度である。出題範囲は前期から後期までの全授業内容を含み、特に以下の領域からの出題が予想される：(1) 記述統計と確率の基礎概念、(2) 推測統計と仮説検定の論理、(3) 分散分析と実験計画法、(4) 回帰分析と相関分析、(5) 一般線形モデル、(6) ベイズ統計学の基礎、(7) 統計ソフトウェア(R と JASP) の操作と結果の解釈。問題形式としては、基本的な用語や概念の理解を問う問題、計算問題、結果の解釈を問う問題、適切な分析手法の選択を問う問題などが含まれる。問題は基礎的な内容から応用的な内容まで幅広く出題され、単なる暗記ではなく、統計的思考力を測定することを目的としている。

学習の総括と重要概念の確認 この試験は、一年間の心理統計学の学習内容を総括する機会でもある。特に重要な概念として、(1) 記述統計と推測統計の区別、(2) 確率分布と標本分布の関係、(3) 仮説検定の論理と p 値の正しい解釈、(4) 効果量と検定力の概念、(5) モデリングの考え方と一般線形モデルの枠組み、(6) 頻度主義統計学とベイズ統計学の違いなどが挙げられる。これらの概念は心理学研究における統計分析の基礎となるものであり、卒業研究や大学院での研究活動においても重要となる。試験問題を通じてこれらの概念を再確認し、心理統計学の全体像を把握することが期待される。

実践的なデータ分析とソフトウェアの利用 本授業では、理論的な内容の学習と並行して、R や JASP などの統計ソフトウェアを用いた実践的なデータ分析のスキルも重視してきた。試験では、ソフトウェアの操作そのものではなく、分析結果の読み方や解釈、適切な分析手法の選択などについての理解を問う問題が含まれる。具体的には、出力結果からの情報抽出、結果の統計学的および実質的意味の解釈、分析結果に基づく結論の導出などが問われる。これらの能力は、心理学研究の実践において必須のスキルであり、試験を通じてこれらの能力が身についているかを確認する。

24.1.3 キーワード

- 心理統計学の総括
- 統計的思考力
- 分析結果の解釈
- 頻度主義統計学とベイズ統計学

24.2 授業情報

■コマの展開方法 授業内試験

■標準シラバスにおける位置づけ 科目番号 5;心理学統計法;総括

24.2.1 予習・復習課題

■予習 本コマは授業内試験であるため、以下の内容について十分な準備をして臨むこと:

1. 前期および後期の授業内容全体を体系的に復習し、重要な概念と用語を理解する
2. 配布資料や教科書の内容を再確認し、特に重要な統計手法とその適用条件について整理する
3. これまでに行った実習や課題の内容を振り返り、分析結果の解釈方法を確認する
4. 試験の詳細(持ち込み可能な資料の範囲など)については、別途配布されるお知らせを参照すること
5. 過去の小テストや課題で理解が不十分だった部分を特に重点的に復習する

■復習 試験終了後は、問題内容を振り返り、不明点があれば次回の授業(総括回)で質問できるよう準備しておくといよい。心理統計学の基礎を確実に身につけることは、今後の心理学研究において重要な基盤となるため、試験を通じて得られた知識の定着を図ることが重要である。また、試験の結果を踏まえて、自分の弱点や苦手分野を特定し、今後の学習計画に活かすことも有益である。

引用文献

- 馬場 真哉 (2020). R 言語で始めるプログラミングとデータ分析 ソシム
- Grimm, Laurence · Yarnold, Paul. (1994). Reading and Understanding Multivariate Statistics. American Psychological Association.
- (グリム, L.G. & ヤーノルド, P.R. 小杉 考司・高田 菜美・山根 嵩史 (訳) (2016). 研究論文を読み解くための多変量解析入門 基礎篇: 重回帰分析からメタ分析まで 北大路書房)
- 橋本 貴充・荘島 宏二郎 (2016). 実験心理学のための統計学——[心理学のための統計学 2]: t 検定と分散分析—— 誠信書房
- Healy, Kieran. (2018). Data Visualization: A Practical Introduction. Princeton Univ Pr.
- (キーラン・ヒーリー 瓜生 真也・江口 哲史・三村 喬生 (訳) (2021). データ分析のためのデータ可視化入門 講談社)
- 川端 一光・荘島 宏二郎 (2014). 心理学のための統計学入門——[心理学のための統計学 1]: ココロのデータ分析—— 誠信書房
- 菊池 聡 (2018). 心理学者は誰の心も見透かせるの?— 学問と偽科学の違い 楠見 孝・日本心理学会 (編) 誠信書房
- 小杉 考司 (2018). 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). R でらくらく心理統計: RStudio 徹底活用 講談社
- Kruschke, John K. (2014). Doing Bayesian Data Analysis. Elsevier.
- (クルシュケ, J.K. 前田 和寛・小杉 考司 (監訳) 前田 和寛・小杉 考司・井関 龍太・井上 和哉・鬼田 崇作・紀ノ定 保礼・国里 愛彦・坂本 次郎・杣取 恵太・高田 菜美・竹林 由武・徳岡 大・難波 修史・西田 若葉・平川 真・福屋 いずみ・武藤 杏里・山根 嵩史・横山 仁史 (訳) (2017). ベイズ統計モデリング: R, JAGS, Stan によるチュートリアル 原著第 2 版 共立出版)
- Lee, M. D. · Wagenmakers, Eric-Jan. (2013). Bayesian Cognitive Modeling: A Practical Course. Cambridge University Press.
- (リー, M.D & ワゲンメーカーズ, E-J. 井関 龍太 (訳) (2017). ベイズ統計で実践モデリング: 認知モデルのトレーニング)
- 松村 優哉・湯谷 啓明・紀ノ定 保礼・前田 和寛 (2021). 改訂 2 版 R ユーザのための RStudio[実践] 入門——tidyverse によるモダンな分析フローの世界—— 技術評論社
- 道又 爾 (2009). 心理学入門一歩手前——「心の科学」のパラドックス—— 勁草書房
- 皆本 晃弥 (2015). スッキリわかる確率統計——定理のくわしい証明つき—— 近代科学社
- 三浦 麻子 (2017). なるほど! 心理学研究法 (心理学ベーシック第 1 巻) 北大路書房
- 長沼 伸一郎 (2016). 経済数学の直観的方法 確率・統計編 講談社
- 中西 大輔・今田 純雄 (編) (2020). あなたの知らない心理学——大学で学ぶ心理学入門—— ナカニシヤ出版
- 大芦 治 (2016). 心理学史 ナカニシヤ出版
- 大久保 街亜・岡田 謙介 (2012). 伝えるための心理統計: 効果量・信頼区間・検定力 勁草書房
- 岡太 彬訓 (2008). データ分析のための線形代数 共立出版

- 清水 裕士・荘島 宏二郎 (2017). 社会心理学のための統計学——[心理学のための統計学 3]: 心理尺度の構成と分析—— 誠信書房
- 清水 裕士 (2021). 心理学統計法 放送大学教育振興会
- 清水 優菜・山本 光 (2022). JASP で今すぐはじめる統計解析入門 講談社
- 下山 晴彦 (2001). 臨床心理学研究の多様性と可能性 下山 晴彦・丹野 義彦 (編) 講座臨床心理学 (pp. 3-24) 東京大学出版会
- Shojima, Kojiro (2022). Test Data Engineering: Latent Rank Analysis, Biclustering, and Bayesian Network. Springer.
- 安井 翔太 (2019). 効果検証入門——正しい比較のための因果推論／計量経済学の基礎—— 技術評論社
- 豊田 秀樹・前田 忠彦・柳井 晴夫 (1992). 原因をさぐる統計学——共分散構造分析入門—— 講談社
- 豊田 秀樹 (2015). 基礎からのベイズ統計学——ハミルトニアンモンテカルロ法による実践的入門—— 朝倉書店
- 豊田 秀樹 (2017). もうひとつの重回帰分析 東京図書
- 山田 剛史・村井 潤一郎 (2004). よくわかる心理統計 ミネルヴァ書房
- 山内 光哉 (2010). 心理・教育のための統計法 サイエンス社
- 野島 一彦・繁榊 算男・山田 剛史 (編) (2019). 心理学統計法 遠見書房
- 結城 浩 (2018). 数学ガールの秘密ノート/行列が描くもの (数学ガールの秘密ノートシリーズ) SB クリエイティブ
- 小杉 考司・紀ノ定 保礼・清水 裕士 (2023). 数値シミュレーションで読み解く統計のしくみ～R でためしてわかる心理統計 技術評論社