

心理学データ解析応用 2/B 詳細シラバス

担当:小杉考司

Last Compiled on 2025.12.3

目次

1	プログラミングの基礎	3
2	プログラミングの実際	4
3	確率関数	5
4	乱数による近似	6
5	一様性, 不偏性, 有効性, サンプルサイズ	8
6	信頼区間	9
7	帰無仮説検定のシミュレーション	10
8	QRPs とサンプルサイズ設計	11
9	確率的プログラミング言語	12
10	モデリングの目から見た検定 1; 二群の平均値差	14
11	モデリングの目から見た検定 2; パラメータの世界とデータの世界	15
12	モデリングの目から見た検定 3; 多群の平均値差を求めるモデル	16
13	混合分布モデル	18
14	確率的プログラミングの応用 1; 項目反応理論	19
15	確率的プログラミングの応用 2; 変化点と折線回帰	20
16	確率的プログラミングの応用 3; 状態空間モデル	22
	Bibliography	23

はじめに

昨今はデータサイエンス、情報科学の領域が非常に隆盛で、コンピュータを使ってデータを分析し、経済の動向や購買行動などの予測に用いられることが広く行われている。

人の行動や考え方をどのようにデータにするかについては、当然ながら心理学には一日の長がある。また、人が頭の中でどのような考え方のプロセスをたどるのか、それをどのように検証するのかについても、心理学はその短い歴史の中で徹底的にその技法を洗練させてきた。このような根源的なレベルでの理論や方法論は時代が変わっても色褪せることなく、また今後ますます必要とされてくる時代になっている。

本講ではデータ解析の応用段階として、より実践的なテーマを扱う。すなわち、**心理尺度が作られる理論的背景と、データの背後のメカニズムを解析する方法**を知ることである。後期配当の心理学データ解析応用 2(旧カリキュラム名心理学データ解析 2B) では、数値計算を通して実践されている推測統計学的手法の原理・メカニズムを知り、ひいてはデータ生成メカニズムを知る技法を学ぶ。

心理学研究では、調査・実験・観察など様々な手法で得られたデータに対し、統計的な処理で「意味があったかどうか」といった判定を下す。データは知りたい全体(知的生命)の一部を使って得たものだから、標本の平均値から母集団の平均値を推測する推測統計学的手法が用いられる。推測統計学では確率や微積分の概念が必要であり、高度で抽象的な数学ツールを実感的に理解することは難しい。畢竟、推測統計のメカニズムの理解がおぼつかないまま、ただ機械から出力される数値を確認するだけになりがちである。生兵法は大怪我の基である。そこで、原理の理解度を上げるために確率変数を数値シミュレーションによって具体的な数値にすることで、抽象度を下げて理解することを考えよう。この方法は、ひいては心理的なメカニズムを数式的に表現することで分析する、**数理モデリング**というアプローチにつながる。このモデリングの基礎的な知識、方法論の習得を目指す。

授業のテーマ

データから意味のある情報を取り出すための、さまざまな分析法を習熟するにあたって、その背後にあって語られることのない「発想」の観点から理解する。数値だけに振り回される状態から脱却し、数値を算出する数式に込められた意味について考える視点を持つ。さらにこれらに習熟することで、どのような研究対象に対してどのような心理統計的アプローチができるかを、俯瞰的に見られるようになる。

一年を通じて伝えたいポイント

乱数生成による確率の理解 ランダムであることを概念的に理解することは難しい。一方で、偶然生じた事象は身の回りに溢れている。確率変数から生じた乱数を用いることで、確率を身近に理解し、推測統計学の理解を深める。

データ生成メカニズム 乱数を用いて、データが生まれてくるメカニズムを考えるトレーニングを行う。データからメカニズムを再現するアプローチは、心理学的事象の深い理解につながる。

統計環境 R と確率的プログラミング言語 Stan による実践 統計環境 R と確率的プログラミング言語 Stan に習熟することで実際に計算し、確認しながら分析を進めることができるようになる。

1 プログラミングの基礎

1.1 授業内容

1.1.1 科目の中でのこのコマの位置づけ

本講義の主眼は「データから意味のある情報を取り出す」ことにある。すなわち、これまでのデータ駆動型 (Data Driven) な発想を逆転し、データがどのように生成されたと考えるかという観点から、データ生成モデル (Data generating Model) による理解を試みる。

モデルは数学的表現がなされ、モデルの形成やデータとの照合 (fitting) には計算機の利用が必須である。初回となる今回の目的は、プログラミングの基本的な考え方を身につけ、次回以降の本格的な運用に備えることである。プログラミングの基礎はコマンドによる命令であり、コマンドの書き間違いはエラーとなって帰ってくる。一見不親切に思えるが、即時反応により学習の効率は良く、コード補完機能などを用いることで簡単なミススペルは回避することができる。小さなものから大きくしていくこと、1 行ずつ実行することなど、基本的な姿勢についても理解する。

また、昨今では chatGPT など生成 AI を用いたプログラミング・アシスト機能が発展しており、積極的な活用をすることで非常に効率よくプログラミングできる。残念ながら、現状では「ハルシネーション (もっともらしい嘘)」の回答が与えられることも少なくない。あくまでも真偽の判断をするのは人間の側に残されており、参考にはできるものの正答が得られるとは限らないことに注意する。

1.1.2 コマ主題細目

パソコンの基礎 スマートフォンやタブレットは「ファイル」を意識しない OS になっており、PC とそれらの境界はますますボーダレス化していくが、プログラミングを行う際はこれらの概念を知っておくことが重要である。PC の基本的な構造 (BIOS/OS/アプリケーション)、ファイルの位置 (パス、作業ディレクトリなど) などについての基礎知識を再確認する。

プログラミングの基礎 プログラミングにあたって重要なことは「思った通りに」動くのではなく、「書いた通りに」動くことである。ミススペルや大文字小文字の違いにも注意が必要である。コードのスペルチェックや補完機能を活用し、また 1 行ずつ確認しながら進めるといったプログラミングに関わる基本的な心構えについて理解する。

いくつかのプログラミング言語 R はプログラミング言語の一種であると言っても良い。プログラミング言語には他にも Python や Basic, C 言語などがある。これらの基本的な関係について理解するとともに、コンパイラとインタプリタという実行形式の違いについても理解する。この違いは今後確率的プログラミング言語を利用する際の知識として生きてくる。

R 言語の基本的な挙動 R および RStudio の最新版を導入し、改めて基本的な活用法を理解する。簡単な四則演算やオブジェクトへの代入、パッケージや関数の利用などを演習的に復習する。tidyverse という概念およびパイプ演算子をはじめとするプログラミング記法についても演習を通じて学ぶ。

生成 AI の効率的な使い方 ChatGPT をコード生成時の補助役として活用する方法について理解する。ChatGPT は人間と違って、些細な質問を繰り返しても決して感情的になることはないため、質問の敷居は低い。一方で、間違えた答えを恥じることもなく、淡々と同じ間違いを重ねることもある。役割や方針を適切に与えた上で、問題を小分割してスモールステップでプログラミングを重ねていくことが基

本である。また、結果が正しいかどうかは、必ず実際に動かしてユーザ側で判断する必要がある。

1.1.3 キーワード

- プログラミング言語
- コンパイラ・インタプリタ
- オブジェクト指向プログラミング
- 生成 AI

1.2 授業情報

■コマの展開方法 講義/遠隔/演習

1.2.1 予習・復習課題

■予習 事前に環境の準備をしておく必要がある。環境の準備についてはいくつかの方策があり、これについては導入資料を参照しながら準備しておくこと。なお、環境準備中に問題が生じた場合はいち早く教員か TA に相談し、実行できるようにしておくこと。

■復習 基本関数の使い方についての演習課題を課す。

2 プログラミングの実際

2.1 授業内容

2.2 授業内容

2.2.1 科目の中でのこのコマの位置づけ

プログラミングの本質は代入、反復、条件分岐である。この3点について理解し、基本的な書き方を学ぶ。実際に R でコードを書きながら、その挙動について確認する。最終的な到達段階として、Fizz-Buzz 問題や行列計算ができるコードを書くこととする。

2.2.2 コマ主題細目

高級言語の基本的な働き 高級言語と呼ばれるプログラミング言語の基本的な働きは、代入、反復、条件分岐である。R では `<-` や `=` で代入を、`for` や `while` で反復を、`if` や `if_else` で条件分岐を行う。これらの表現は言語間を通じて共通なものが多いため、その基本的な振る舞いを確認することは技術の一般化に役立つ。

オブジェクトの型 R は変数を宣言せずに利用できるところが初心者向けの長所ではあるが、R の中にストックされるオブジェクトの種類や型について正しく理解しておくことが、今後のデバッグの上で役にたつ。変数としての数値型 (Int, Double, Complex), 文字型 (Char, Factor), 論理型 (logical) や、オブジェクトのベクトル、行列、配列、リスト、データフレームについて理解する。

関数を作る コマンド (文) は連なってスクリプトとなる。反復されるスクリプトは関数化すると効率的である。ここでは R を用いた関数の定義の仕方について学ぶ。引数やデフォルトの値について知ることは、へ

ルプを見る時の役にも立つ。また関数化はスパゲティ・コードになることを避ける、カプセル化された表記方法への第一歩であり、オブジェクト指向プログラミングへの第一歩でもある。

2.2.3 キーワード

- オブジェクトの型
- 代入
- 反復
- 条件分岐
- 関数

2.3 授業情報

■コマの展開方法 講義/遠隔/演習

2.3.1 予習・復習課題

■予習 RStudio におけるプロジェクト単位での管理など, R/RStudio の基本的な使い方については以後も逐一支持することはないので, これまでのことをしっかり復習し, 本講義の予習とする。

■復習 反復計算の練習課題, 条件分岐の練習課題など, 複数の課題にしっかり取り組むこと。

3 確率関数

3.1 授業内容

3.1.1 科目の中でのこのコマの位置づけ

心理統計は個人差や偶然性を表すために確率の考え方が必須である。しかし確率の概念は数学的に高度かつ抽象的であり, 直感的に理解することが難しい。概念的には, 空間を標準化したときの相対的な大きさを指標化したもの, すなわち面積の一種と考えるとわかりやすいし, 複雑な性質よりも最低限守るべき公理の方からアプローチした方が理解しやすい。ただしこのとき, 確率変数とその実現値の違いに注意する。とはいえこのコースの目的は, 数学的内容を座学で学ぶのではなく, あくまでもプログラミングによる演習を通じて具体的に理解することにある。本講ではまず概念的な理解の復習からすすめる。なお以後, kosugi2023 を副読本としながら授業を進めていく。

3.1.2 コマ主題細目

確率の基礎 高校数学までの範囲で学ぶ確率は, すべての場合の数を全体としたうち, 該当する事象の生じる回数を相対頻度で表すものであった。一方我々は, 天気予報で降水確率が 30% というように, 確率的な表現を日常的にも用いる。天気予報は未来の予測であるから, 「起こりうるすべての未来」を書き出して生じる天気のを割合を表現できるはずもない。こうした不確実なものであっても, 確率という言葉で表現できるようにコルモゴロフが公理として最も基本的な枠組みを提供した。ここでは認識論に深く立ち入ることはせず, 確率という数字が従うべき法則について確認する。

→ Maeda2017 の P.78–83, kosugi2023 の P.61–62.

R の確率関数 R には確率関数がいくつか用意されており, 接頭語 `d, p, q, r` と関数名からなる。ここでは正規分布を例に, `d, p, q-norm` がそれぞれ何を表しているかを確認する。確率は確率密度関数で表される領域の面積であり, 積分によって計算されることを近似計算で理解する。

→kosugi2023 の P.66–77.

確率変数の期待値と分散 確率変数の期待値 (加重平均) と分散について, 定義式を確認する。また正規分布においては期待値が位置パラメータ μ に, 分散がスケールパラメータ σ^2 に一致することを確認する。

→kosugi2023 の P.83–86.

3.1.3 キーワード

- 確率変数と確率変数の実現値
- `dnorm, pnorm, qnorm, rnorm`
- 確率分布関数

3.2 授業情報

■コマの展開方法 講義/遠隔/演習

3.2.1 予習・復習課題

■予習 一年時の「データ解析基礎」ですでに確率については学んでいる。テキストや関連資料を見て改めて確率の概念を復習しておくことが, このセクションの予習になる。

■復習 R の正規分布を例に, 関数の形とパラメータの関係, 期待値と分散について学んできたが, 正規分布以外の確率分布関数についても定義や性質を調べつつ, R の関数をつかって描画・計算して理解を進めておくと良い。

4 乱数による近似

4.1 授業内容

4.1.1 科目の中でのこのコマの位置づけ

確率分布についての数学的位置付け, 理解について確認したところで, より具体的かつ体験的に理解するために, 乱数を用いることで近似できることを学ぶ。計算機である以上, 完全な乱数ではなく疑似乱数に過ぎないが, 人間の人為的な乱数生成よりも, より「ランダムな」数字を作ることができる。さらに計算機の乱数はすでに確率分布に従うものが準備されているから, 試しながら理解を深めることができる。また今後のベイズ推定に向けて, 名も無分布からの乱数であっても同様のことができることを理解する。

4.1.2 コマ主題細目

計算機の乱数 計算機の擬似乱数は、所謂「ガチャ」のような振る舞いをする数字のことである。一様乱数を例にサイコロの出目を再現するプログラムを通じて、擬似乱数の特徴を学ぶ。とくに乱数の種 (seed) を固定することで、再現可能であることにも注意する。

乱数による形状の近似 ここまでは確率密度関数を使っただけの話であったが、乱数を使うことで確率変数の実現値を使った近似が可能である。正規乱数を例に、まずは乱数の数を増やしながら、ヒストグラムの形状が徐々に正規分布に近づいていくことを確認する。また、一様乱数など他の確率分布の例でも確認してみる。

乱数による値の近似 p, q -norm を駆使して計算した面積の問題も、乱数を使って近似することができる。乱数は確率変数の実現値であり、R のもつ統計関数を使った集計が、積分などの計算に対応するものであることがわかる。期待値と分散の計算を乱数によって近似する。近似の精度は生成する乱数の数に依存することを理解する。

→kosugi2023 の P.83–86.

事後分布からの乱数に向けて 正規乱数や一様乱数、そのほかにも F, t, χ^2 分布など、聞いたことのある確率分布であれば乱数を生成して計算することで、数学的 (解析的) アプローチを避けることができる。また、特殊な例に思えるかもしれないが、とくに有名な名前のついていない確率分布であったとしても、そこから乱数が生成できれば、ここまで学んできたことを利用して近似計算が可能なのはである。この後のベイズ統計学的アプローチの理解に向けて、この論点を先に押さえておきたい。

4.1.3 キーワード

- 疑似乱数
- 乱数の種
- 正規乱数, 一様乱数
- 事後分布

4.2 授業情報

■コマの展開方法 講義/遠隔/演習

4.2.1 予習・復習課題

■予習 前時の R の確率関数の使い方を十分に復習しておくこと。とくに積分の計算によって確率分布の一部の面積が算出できることを理解する。

■復習 kosugi2023 のコラム P.87–90 にある正規分布の再生性について、テキストに沿って確認しておく。確率分布同士の計算も、解析的にではなく近似的に分析することができることを知ると、自信につながる。

5 一致性, 不偏性, 有効性, サンプルサイズ

5.1 授業内容

5.1.1 科目の中でのこのコマの位置づけ

ここからは推測統計学における確率分布の利用について理解する。すでに推測統計の基礎的な発想は一年時に学んでいるが, 結果を追従するための理解ではなく, 仕組みを理解するための一歩進んだ考え方を学ぶ。改めて母集団と標本の関係について理解をした上で, 確率分布や数理統計の技術がどのように使われているかを理解する。kosugi2023 の第 4 章を参考にする。なお, 有効性の理解についてはテキストを参照するにとどめる。

5.1.2 コマ主題細目

推測統計の基礎 母集団, 標本という関係, また母数と標本統計量の区別を改めて確認する。R の `sample` 関数を使って, 全体から一部を取り出す練習をすることでイメージがしやすくなるだろう。また標本統計量は標本を取るたびに変わりうる値, すなわち確率変数であるから, 標本統計量の従う分布すなわち標本分布についても理解する必要があることを確認する。

→kosugi2023 の P.119–120.

一致性の理解 標本から母数を推測するとき用いられる標本統計量を推定量というが, この推定量が持つべき望ましい 3 つの性質について, シミュレーションを通じて学ぶ。まず一致性について, サンプルサイズが大きくなればなるほど母数に近づく性質である。正規乱数のサンプルサイズを変えてこれを確認する。

→kosugi2023 の P.121–128.

不偏性の理解 標本分散と不偏分散の違いは, 学んだ当初はその区別がわかりにくかったかもしれない。ここでシミュレーションを通じてその理解を深める好機である。不偏性は推定量の期待値が母数に一致することであるから, 正規乱数の標本分散と不偏分散の計算を反復し, その平均を取ることでいずれが推定量として適切であるかを確認する。

→kosugi2023 の P.128–137.

サンプルサイズと中心極限定理 ここまで正規乱数の例が多かったが, 中心極限定理によると母集団が正規分布でなくてもその標本平均の分布は正規分布に近づく。t 分布を例にこのことを確認すると同時に, 心理学実験などサンプルサイズが小さい場合にどのような注意が必要かについて学ぶ。

→kosugi2023 の P.140–144.

5.1.3 キーワード

- 母集団と標本
- 一致性

- 不偏性
- 有効性
- 中心極限定理

5.2 授業情報

■コマの展開方法 講義/遠隔/演習

5.2.1 予習・復習課題

■予習 一年時の「データ解析基礎」のテキストを復習しておくこと。第 16 講 (後期第一回目) が本時に対応する。

■復習 テキストには正規分布以外の確率分布を使って例が多く載っている。参考にしながら一致性・不偏性について理解を深め、また有効性についても一読しておくこと。

6 信頼区間

6.1 授業内容

6.1.1 科目の中でのこのコマの位置づけ

標本統計量が確率変数であるということは、心理学実験や調査などで得られるデータとその統計量をそのまま一般化できないことにつながる。正規母集団から得られた標本平均は幸いにして、母平均の推定量として好ましい性質を有しているが、それにしても標本平均すなわち母平均と考えるのは極端であると言わざるを得ない。標本統計量が確率変数であり、確率分布 (標本分布) の性質がわかっているのなら、それを使って区間推定することが考えられる。ここでは区間推定についてもシミュレーションを通じて理解する。

6.1.2 コマ主題細目

点推定と区間推定 点推定と区間推定の区別を確認した上で、シミュレーションによって点推定が母数と一致する割合、区間推定が母数を含む割合を確認する。前者がゼロであることから、区間推定の必要性が理解でき、かつどのように区間を設定すれば良いかという問題に気づく。

→kosugi2023 の P.145–146

母分散が既知の場合 すでに「データ解析基礎」において標準正規分布を利用した区間推定について、95% 区間に対応する ± 1.96 という数字を使った例を学んだ。今回はシミュレーションを通じて、標本平均 $\pm 1.96 \frac{\sigma}{\sqrt{n}}$ が正しく母平均を含んでいる回数をチェックし、このことを確認する。

母分散が未知の場合 同じく、母分散が未知の場合は t 分布を使って区間推定をするのであった。ここでもシミュレーションを通じて、区間推定が正しく母平均を含んでいる回数をチェックし、このことを確認する。可視化することでこのことが理解を深める。

→kosugi2023 の P.146–156.

6.1.3 キーワード

- 点推定
- 区間推定
- t 分布

6.2 授業情報

■コマの展開方法 講義/遠隔/演習

6.2.1 予習・復習課題

■予習 一年時の「データ解析基礎」のテキストを復習しておくこと。第 17 講 (後期第二回目) が本時に対応する。

■復習 t 分布の数式について、テキストのコラムを一読しておく。余裕があれば、テキスト 4 章後半の相関係数の信頼区間などにも目を向けるとよい。

7 帰無仮説検定のシミュレーション

7.1 授業内容

7.1.1 科目の中でのこのコマの位置づけ

標本統計量が確率変数であるなら、「二群の差」といった標本の群間差も確率的に変動するし、これに意味があるかないかといった判断も確率的になる。この判断が誤ったものになる水準をコントロールしようというのが帰無仮説検定の考え方である。このことに注意しておけば、帰無仮説検定の手続きも「米印を探す単調作業」でないことは理解できるだろう。改めて帰無仮説検定の手順を確認するとともに、t 検定とシミュレーションを通じて検定の精度と意義の理解を深める。なお kosugi2023 の第 5 章を参考にする。

7.1.2 コマ主題細目

帰無仮説検定の論理 帰無仮説と対立仮説、判断の手続などを再確認する。また、用語としてのタイプ 1 エラー、タイプ 2 エラー、検定力、および帰無仮説が従う分布 (帰無分布) についても解説を加える。

→kosugi2023 の P.187–191.

R による検定の実際 二群の平均値差の検定を例にとる。データを乱数から生成し、どのような手順で検定が行われているか、そのアルゴリズムを復習する。

→kosugi2023 の P.191–196.

タイプ 1 エラー確率のシミュレーション 心理学の研究実践において、タイプ 1 エラー、タイプ 2 エラーがどれぐらいだったか、検定力がどれぐらいあったかを知る術はない。しかしシミュレーションすることによって、何がどの程度起こりうるかを体験することができる。これをもって、自らの研究実践に対する統計的指標の意義を考える一助にしてもらいたい。また、t 検定においては Welch の方法がデフォルト

トで用いられるが、等分散性の仮定が成立しない場合はどのような問題が起きるかも確認できる。

→kosugi2023 の P.196–203.

7.1.3 キーワード

- 帰無仮説と対立仮説
- 帰無分布
- サンプルサイズ, タイプ 1 エラー, タイプ 2 エラー, 検定力
- 等分散性の仮定

7.2 授業情報

■コマの展開方法 講義/遠隔/演習

7.2.1 予習・復習課題

■予習 一年時の「データ解析基礎」のテキストを復習しておくこと。第 18–20 講が本時に対応する。

■復習 サンプルサイズや二群それぞれの分散の値をさまざまに変化させて、シミュレーションの結果がどのように変わるかいろいろ試してみよう。

8 QRP とサンプルサイズ設計

8.1 授業内容

8.1.1 科目の中でのこのコマの位置づけ

ここでは帰無仮説検定の仮定を逸脱した場合に何が起こるか、悪しき研究実践の例を仮想的に体験することでその問題を理解する。帰無仮説検定はデータに対して厳格なルールを決めてから実践すべきもので、いわばスポーツのゲームのような側面がある。ラフプレーをするとゲームが成立しないように、科学におけるラフプレーは真実を見誤るという問題を引き起こす。とくにサンプルサイズや帰無仮説を事後的に変えることで、本来コントロールされているはずのものがそうならなくなってしまう。これらについて、kosugi2023 の第 6 章を参考に話を進める。

8.1.2 コマ主題細目

心理学の再現性問題 2000 年ごろから指摘され始めた、心理学の再現性問題を紹介する。その一因として挙げられる心理統計法の誤用と QRP(問題ある研究実践) の存在を知り、このことが科学としての心理学の根幹を揺るがす大問題であることを共有したい。資料として Ikeda2016 を用いる。

QRP の実践 QRP を仮想空間上で実践して、どのような問題が生じるかを見ていく。サンプルサイズを徐々に増やしていく(「もう少し頑張ってデータを取ろう!」)ことがなぜ悪いのか、有意なところだけ報告する(帰無仮説の事後的な変更)などによって、タイプ 1 エラーの確率が制御できない様子を確認する。

→kosugi2023 の P.6–8 および P.225–232.

非心分布をつかったサンプルサイズ設計 以上のことから、帰無仮説検定を行う際は事前にすべての設定・ルールを定めておく必要があることがわかる。実践者が恣意的に決められるのはサンプルサイズであるから、サンプルサイズの設計が重要になってくる。ここで対立仮説が従う分布として非心分布を導入し、これに基づくサンプルサイズ設計の例をシミュレーションで確認する。

→kosugi2023 の P.6–8 および P.238–224.

非心分布をつかわないサンプルサイズ設計 二群の平均値差の検定のような、比較的簡単な実験計画においてはサンプルサイズ設計の理論も単純なほうである。とはいえ、非心分布など慣れない確率分布の利用に戸惑うこともあろう。そこで多少力技的ではあるが、乱数生成を繰り返すことで実際の程度の差・効果がどの程度の割合 (確率) で検出できるかを検証するシミュレーションを行う。この方法は非心分布がわからないときや、複雑な実験計画にであっても応用が可能である。

8.1.3 キーワード

- 再現性問題
- QRPs
- 非心分布
- 例数設計

8.2 授業情報

■コマの展開方法 講義/遠隔/演習

8.2.1 予習・復習課題

■予習 一年時の「データ解析基礎」のテキストを復習しておくこと。第 19 講が本時に対応する。

■復習 分散分析の例数設計については、テキストを参考に自ら実践することが望ましい。

9 確率的プログラミング言語

9.1 授業内容

9.1.1 科目の中でのこのコマの位置づけ

これまでは R にビルトインされている確率関数を使って、シミュレーションや検定のロジックを確認してきた。ここまでは、確率を反復試行におけるある実現値の出現割合と考えてきたことになる。この考えに基づく確率の定義は頻度主義的ともいわれるが、同じ確率を使った推論でも他の方法が存在する。ベイズ的確率の考えがそれで、この場合の推論は事前分布と確率モデルから導出される事後分布を活用して行われる。心理統計でよく用いられるモデルは、一様分布を事前分布とし、確率モデルは正規分布を仮定することがほとんどであり、事後分布も解析的に導出することができるが、ここではより一般的な確率モデルに応用することを考える。確率的プログラミング言語を持ちることで、尤度と事前分布を設定することで事後分布発生器を作ること

ができ、いかなる事後分布からでも乱数を生成することができる。乱数によって確率分布が近似できることはすでに学んだとおりである。ここでは専門の確率的プログラミング言語 Stan を導入し、以後の学習に備える。

9.1.2 コマ主題細目

ベイズ推論の基礎 推定法の一種、ベイズ推定はベイズの定理に基づいている。改めてベイズの定理を解説し、条件付き確率や尤度、事前分布、事後分布といった用語についても改めて確認する。

→Maeda2017 の Pp104–123 ほかに枚挙に遑がない

事後分布の生成 マルコフ連鎖モンテカルロ法は、事後分布の生成と乱数の生成が合体した計算技術である。事後分布の関数の形はわからなくとも、そこから乱数を生成することで形状を近似し、分析することができる。これらの基本的な仕組みを理解する。

→Maeda2017 の Pp.147–194

確率的プログラミング言語 Stan マルコフ連鎖モンテカルロ法の演習にあたって、確率的プログラミング言語 Stan を導入する。Stan は変数の宣言が必要であること、ブロックに分割されていること、行の終わりにセミコロンを入れることなど、R 言語よりも C 言語に近い書き方が必要である。RStudio には Stan ファイルのサンプルも入っているので、これらを活用して簡単なコードを書いてみる。

→Maeda2017 の Pp.407–425.

事後分布による解析 Stan を R から利用すると、結果として得られるオブジェクトに非常に多くの情報を含んでいることになる。ここでは cmdstanr の出力結果を確認し、またコンパクトに出力をまとめる関数を作りながら、事後分布をどのように可視化、記述するかについて学ぶ。

9.1.3 キーワード

- ベイズ推定
- 事前分布, 事後分布
- 確率的プログラミング言語
- Stan

9.2 授業情報

■コマの展開方法 講義/遠隔/演習

9.2.1 予習・復習課題

■**予習** Stan の導入は大学の PC 教室を利用するが、そうでなければローカルに環境を構築する必要がある。ローカル環境での利用を希望するものは、事前に調べて用意しておく。

■**復習** Stan のサンプルコード、可視化の関数やパッケージなど、演習の分量が多いため、時間内に終わらなかったものは資料を参考に必ずフォローアップしておくこと。

10 モデリングの目から見た検定 1; 二群の平均値差

10.1 授業内容

10.1.1 科目の中でこのコマの位置づけ

データ生成モデリングの観点を踏まえた上で、検定的アプローチとベイズ的アプローチの違いを学ぶ。

二群の平均値の差の検定、いわゆる対応のない t 検定の場合は、同一の正規分布から得られたデータに対して平均値の差があると判断して良いかどうかを判断するという枠組みであった。これらの前提と判断基準を確認し、それがデータ生成モデルの観点ではどのように表されるかを検証する。ここで結果が分布として推定されること、差があるかないかというのは二群の推定された平均値の差であることから、生成量を使って平均値の差を出力することを考える。ここで効果量に改めて目を向けるとその理解が進む。

10.1.2 コマ主題細目

t 検定の仮定 二群の平均値の差を検定するときは t 検定が利用されるが、データが正規分布から得られているという仮定、分散が同じであるという仮定などを踏まえて設計図を書き、これを Stan で表現することを考える。あらためて t 検定のやり方や結果と比べてみることで、モデリングがデータ生成メカニズムに注目していること、パラメータの推定を行なっていることなどが確認できるだろう。また等分散性の仮定を外す方法についてもすぐに応用ができる。

帰無仮説検定と二群の差の検定について、→Yamada2004 や一年次の資料をもとに確認しておく。

差の分布 検定は推測に加えて判断を行っていた、ということを改めて確認するとともに、帰無仮説検定では母平均の差をターゲットにしていたことを確認する。MCMC は母集団からの代表値であるので、推定された結果を使って差を表現することができる。これは R 側で得られたサンプルで行っても良いし、Stan の生成量を使っても良い。ここで generated quantities ブロックの考え方を導入し、平均値の差の分布を確認すること、帰無仮説検定が差の分布の一点についての仮説であったことを確認する。また一方が他方よりも大きくなる確率はどれぐらいかとか、一方と他方が c 以上に違っている確率はどれぐらいか、と言ったことが生成量を使って計算することができるようともいえる。

帰無仮説検定を省みる ここまでくると、帰無仮説検定のロジックや考え方について別の視点から見ることができるようになるであろう。まずは帰無仮説と対立仮説という対立のさせ方の不平等さである。帰無仮説は一点についての仮説であり、対立仮説はそれ以外であればなんでも良い、という非対称な関係になっていた。それを省みると、差があるかないかといった二値判断に陥ることがいかに危険であるかわかるだろう。また量的な判断ができないことから、効果量を合わせて報告することが望ましいとされている。効果量とは、標準化された差の大きさのことであり、生成量を使って簡単に算出することができる。また方向性を持った検定について、片側・両側検定などで考えられてきたが、生成料を使えば自然にそれが検証できることがわかる。ただしこれらの検証の仕方は、今回のデータと仮定されたモデルという前提の上で成立する程度であって、過度な一般化にはならないように注意する必要がある。

10.1.3 キーワード

- t 検定

- 生成量
- 効果量
- 片側検定, 両側検定

10.2 授業情報

■コマの展開方法 講義/遠隔/演習

10.2.1 予習・復習課題

■予習 Stan の基本的なブロック構成, 設計図からコードに落とすやり方を確認しておこう。

■復習 データのサイズが変わるだけでなく, 平均値の差, 効果量が変化したときの t 検定の結果とベイズ推定の結果がどのように変わるのか, さまざまなケースを想定して「遊んで」みるとよい。加えて, そのほかの仮説検定がどのようなデータ生成メカニズムで表現できるかを考えることは, 次回以降の準備にもつながる。

11 モデリングの目から見た検定 2; パラメータの世界とデータの世界

11.1 授業内容

11.1.1 科目の中でこのコマの位置づけ

データ生成メカニズムの観点から帰無仮説検定を省みた場合, その仮定やメカニズムを明記する必要があることから, 拙速な解釈に陥る危険性を避けることができる。また事後予測分布を作ることで, 柔軟な仮説を考えられることなども示された。

本時は, 同じく事後予測分布を使いながら, パラメータでなくデータのレベルでの比較ができることに言及し, 実質的に差があるとはどういうことであるかを考える。パラメータの世界, データの世界を分けて考えられるように注意を促す。まずは事後予測分布をみることで, モデルが現在のデータを正しく再現しているかを見ることで, 視覚的にモデルの正しさが検証できることを確認する。その上で, 新しく作られた分布の特徴から, 閾上率や優越率などを計算することができる。これらはデータに基づく予測であるから, より具体的で実感をしやすい予測として使えるだろう。翻って, 仮想データを生成して検証する, パラメータリカバリの手法を学ぶ。この方法では真値やサンプルサイズを自由に設定し検証できることから, 例数設計に応用することが可能である。ベイズ推定をしない場合であっても, シミュレーションによる例数設計が有用であることを理解する。

11.1.2 コマ主題細目

事後予測分布 推定値をつかって, 新たにデータを生成した場合どのようなことが言えるか。事後予測分布を描くことでモデルの正しさが確認できる。事前分布の特徴を反映した, 事前予測分布についても触れる。

事前と事後の予測については →Lee201709 の Pp.38–42 が詳しい。

データレベルの仮説 これまで考えてこられた仮説は, パラメータについての仮説であった。一方, 事後予測分布が新しいデータを作っているのであれば, そこからデータレベルの仮説を考えることもできる。閾上率や優越率といった, データのレベルでの仮説を検証したり検討したりすることを考える。これら

の視点は帰無仮説検定およびモーメント法による算出よりもわかりやすいかもしれないし、統計的に差があるということがどの程度意味のあることなのかを実感するのに役立つだろう。

優越率、閾上率については、→Toyoda201606, Pp.69–70.

パラメータ・リカバリ 事後予測分布は乱数発生による新しいデータの生成である。であれば乱数発生のアプローチは、理論に従う仮のデータを生成することができる、ということでもある。シミュレーションとして仮にデータを作ってみて、サンプルサイズがどの程度であればどの程度正確な推定ができるのか、と言った理論的な検証をすることができる。これはパラメータ・リカバリという試みでもあり、モデルが複雑になって行ったときに正しく機能するかどうかをチェックする方法でもある。また、サンプルサイズも自由に変えることができるのだから、どの程度のサンプルがあればどのような結果が得られるのか、と言ったシミュレーション、あるいは実験前のサンプルサイズ設計にもつながる。

モデリングの基礎的手順について、→Matsuura201610 の Pp.12–81.

11.1.3 キーワード

- 生成量
- パラメータ・リカバリ
- 事後予測分布
- 例数設計

11.2 授業情報

■コマの展開方法 講義/遠隔/演習

11.2.1 予習・復習課題

■予習 generated quantities ブロックの書き方について、数値を色々変えて確認しておくといよい。

■復習 さまざまなサンプルサイズ、効果量の仮想データを生成し、帰無仮説検定やベイズ法による推定を繰り返すことで、各手法の長所や短所を考えることができる。遊び心を持って、さまざまな状況生成して、実際に試してみること。

12 モデリングの目から見た検定 3 ; 多群の平均値差を求めるモデル

12.1 授業内容

12.1.1 科目の中でこのコマの位置づけ

ここまで generated quantities ブロックの活用で、事後分布、事後予測分布を生成できること、パラメータの世界、データの世界それぞれでの仮説が検証できることを見てきた。またパラメータリカバリの方法を見ることで、データ生成モデルを分析前に活用する方法についても見た。

続いて多群の平均値差を求めるモデルを考える。まずは一要因 3 水準の Between モデルから、3 つの平均値をバラバラに求めること、生成量から差分を計算することを考える。データやパラメータレベルでの仮説

的な検証方法を再確認し、帰無仮説検定の枠組みで考えなければならなかったアルファ水準のインフレ問題が生じないことを、確率の考え方の違いに即して理解する。続いて差をモデルに組み込む方法を考える。ここでパラメータの数の制約をかける方法として `transformed parameters` ブロックの使い方を導入する。またパラメータの数の制約は、自由度の概念と深く関係していることへの洞察を得る。またパラメータリカバリのコードがリバースエンジニアリングのコードと同じもので、鏡合わせの関係にあることを確認する。続いて交互作用が含まれるモデルを考える。ここで制約からどれだけのパラメータが必要か、どのように組み上げるかを学ぶことができる。

12.1.2 コマ主題細目

要因計画モデル 一要因 3 水準 Between モデルを考える。対応のない二群の時のように、これは素直に 3 群のモデルとして表現できるし、群間の差を生成量として計算できることを再確認する。また検定と違って差の大きさを直接検証すること、確率的判断を含まないことから、アルファ水準のインフレ問題に悩む必要がないことがわかる。この考え方は、確率の捉え方の違いにもつながることに留意する。

パラメータの変形と制約 3 群のうち、ある群を基準にした差分を直接パラメータとして推定するモデルを考えると、2 つの差分を計算することができる。またある群を基準におかなくとも、全体平均を基準に置くことができるが、その場合は差分のパラメータに制約をかける必要がある。これらの制約を含んだモデルを、`transformed parameters` ブロックを使って作ることを確認する。ここでパラメータの自由度について理解する。

モデルの洗練 技術的な問題であるが、多群モデルの場合は一般的に描画するためにも、整然データの形式に整えておくことが望ましい。書いたモデルの一般化という観点から、変数にできるところは変数にするなど、コードの洗練を試みる。ここで群を識別する変数を導入することは、今後の個人内反復測定モデルにも応用できる点であるので、しっかり理解する。

パラメータリカバリ これらのモデルのパラメータリカバリから、仮想データはデータ生成モデルを逆転させるだけで出来上がることがわかり、要因計画を裏側から眺めるような、新しい観点からの理解が進むと考えられる。X

12.1.3 キーワード

- 要因計画
- 整然データ
- 生成量
- パラメータリカバリ

12.2 授業情報

■コマの展開方法 講義/遠隔/演習

12.2.1 予習・復習課題

■予習 パラメータリカバリの必要性など、データ生成モデルのアプローチにおける標準的な手順を再確認する。また対応のない二群の平均値差の検定、要因計画の検定についてこれまでの復習をしておくことで、今

回の内容の理解が深まるだろう。

■復習 要因数が増えた場合どのようなになるか、またその都度 Stan モデルを書き換えなくても良くなるような一般的な書き方について、自分なりに試行錯誤することが望ましい。

13 混合分布モデル

13.1 授業内容

13.1.1 科目の中でのこのコマの位置づけ

ここまで GLMM, HLM とさまざまな分布を組み合わせて利用するモデルについてみてきた。

今回は混合正規分布モデルと 0 過剰ポアソンモデルを考える。これまではデータが全体的に均質であることを想定していたが、そもそも異なる種類のデータが混じり合っていると考えられる場合、異なる分布をあてがう方がよい。ここでどちらの分布に属するかがある種の確率で代わり、それに従って続くモデルが異なるという、離散確率分布による条件分岐をデータ生成メカニズムに導入することになる。またこれを Stan で実行する場合には、離散変数が直接扱えないことから、すべての可能な場合を数え上げて足し合わせるという周辺化して消去するというテクニックを利用することになる。技術的に高度な側面もあるが、条件分岐をデータ生成メカニズムに組み合わせることができれば表現力は一気に広がることになる。

13.1.2 コマ主題細目

混合分布モデル 混合分布モデルは確率的なクラスター分析 (Model Based Clustering) でもある。階層線形モデルと違って、クラスター分析はデータの背後にあるクラスが明示的に示されておらず、データの適合から考えることになる。まずデータの可視化によってまずその可能性に気づき、これをどのようにモデル化できるか、そのアイデアを理解する。具体的には、ベルヌーイ分布など離散変数のパラメータによって条件分岐が発生し、各条件のもとで確率モデルが描かれることになるが、この確率モデルをどのように統合するかということを設計図の段階で理解する。設計図での理解を踏まえて、Stan での実装レベルでの理解に進むことができるのだから。

→Matsuura201610 の Pp.209–213.

周辺化消去 Stan は離散確率分布を直接モデルの中に組み込むことはできない。そこで `log_sum_exp` という特殊な関数を使って、すべての場合わけを行った離散モデルの統合を考える。まずターゲット記法について理解し、続いて周辺化消去の書き方を理解する。また混合正規分布モデルの場合、ラベルスイッチングの問題が生じることが考えられるから、`ordered vector` の型を利用することが多い。こうした特殊な関数やベクトルについても理解を深める

→Matsuura201610 の Pp.203–208.

0 過剰ポアソン データの性質を考えると、必ずしも正規分布モデルばかりではなく、離散変数など一般的なモデルまで拡張することが可能である。そこで野球データなど具体的な分布情報とともに、ゼロ過剰ポアソン分布を使ったモデルを考える。データ生成のメカニズムがより具体的、実践的に考えることができるので、モデリングによる分析の自由度が高まることを実感できるだろう。

→Matsuura201610 の Pp.82, 83, 85–89.

13.1.3 キーワード

- クラスター分析
- 混合分布モデル
- 周辺化消去
- ゼロ過剰ポアソン分布

13.2 授業情報

■コマの展開方法 講義/遠隔可/演習

13.2.1 予習・復習課題

■予習 データの描画から気づかされることは無数にある。探索的にデータプロットができるように、`ggplot` やデータハンドリング技術について復習しておくことが望ましい。

■復習 混合分布モデルが応用できるようなデータ例を身の回りから考えてみるとよい。その上で、データがどのような背景から生成されているかの設計図を書き、設計図からコードに起こすという手順を一步步確認しておこう。

14 確率的プログラミングの応用 1; 項目反応理論

14.1 授業内容

14.1.1 科目の中でこのコマの位置づけ

ここまでで、ベジアンモデリング・アプローチによって従来の統計モデルがどのように変わるかを見てきた。以後はアラカルト的に、確率的プログラミングの応用による柔軟なモデリング例のトピックスを取り上げる。最初に扱うのは、項目反応理論のモデリングである。尺度作成の文脈で、理論的概要は心理学データ解析応用 1 のシラバスやテキストを参考にしてもらいたいが、改めて確認するとともに確率的モデリングとして実装する。

確率モデルとして考えると、0/1 の反応に対するロジスティック回帰の応用であり、実装自体は既有知識の応用で可能であろう。コーディングのポイントとして、`long` 型データ (`tidy data`) にしておくことで欠測値が含まれる場合も対応できるようになることが挙げられる。

14.1.2 コマ主題細目

ロジスティックモデルの復習 項目反応理論を R で実行する場合は、`ltm` パッケージなど専用の関数群がすでに存在する。しかしここでは 0/1 の二値反応に対する確率分布である、ベルヌーイ分布のパラメータにモデルを組み込む形として、フルスクラッチでコードを書き直す必要がある。あらためてロジスティックモデルの概略を説明し、テストデータとの関係について理解する。

ロジスティック回帰モデルでの実装 ロジスティック回帰分析を拡張した 1PL, 2PL, 3PL モデルそれぞれ

を, `transformed parameters` ブロックで記述することを演習で学ぶ。

整然データでの分析 データを整然データの形にして分析することで、欠測値が含まれないデータセットを作って分析に応用することができる。ここでは個人と項目それぞれを識別する変数が必要になるが、これまで学んできた技術で十分対応可能であると考えられる。

14.1.3 キーワード

- 項目反応理論
- 1PL ロジスティックモデル
- 2PL ロジスティックモデル
- 3PL ロジスティックモデル
- 整然データ

14.2 授業情報

■コマの展開方法 講義/遠隔可/演習

14.2.1 予習・復習課題

■**予習** これまでの知識や技術を組み合わせて問題に対応することになる。項目反応理論とロジスティックモデル、GLM におけるロジスティック回帰分析、データハンドリングにおける整然データの考え方など、これまでの資料に戻って復習しておくといい。

■**復習** 自分で描いたモデルが R のパッケージが出す答えとどの程度一致するのかを確認しておこう。また欠測値がある被験者の被験者母数は、その確信区間が広がると考えられる。なぜそうなるかを改めて考え、実際のデータ適用例で確認しておこう。

15 確率的プログラミングの応用 2; 変化点と折線回帰

15.1 授業内容

15.1.1 科目の中でこのコマの位置づけ

15.1.2 コマ主題細目

変化点検出は、時系列的なデータの中に異なる 2 つの平均値を持つ群があることをモデリングする手法である。とくにある時点から異なる群に属するという系列的な意味があることと、変化点があるとすればどのあたりになるかという「変化点の位置的不明確さ」を確率分布で表現し、データから検出するという観点は、確率モデルの表現の自由さとデータとの接合を許す確率的プログラミング言語の面白さを味わうには最良の材料である。

まずは混合分布モデルのように、2 つの群を分類するモデルを再確認し、その上で時系列的なデータという既有知識から「変化点」という考え方の導入、モデリングへと繋げる。またデータによっては、一定の点を期に線形モデルの傾きが変わるような表現が可能なのがある。この変化点と回帰分析を融合させた、折線回帰モデルを考えることで、固定的なモデルを超えた柔軟なモデリングが可能であることを理解する。

ただしここで使うデータは時系列的なものであるから、一般的な回帰分析の前提であるサンプルの独立性がない。その意味で不適切なモデルであることに注意し、次の時系列分析へと繋げる。

混合分布モデル データは可視化することが重要であり、見れば明らかに異なる状態の混合であることがわかる場合がある。具体例とともに可視化を行い、またこれまで学んだ混合分布モデルで表現できることを再確認する。ここで用いるデータは、小杉の体重記録データを用いる。

変化点検出 データの横軸が時系列的な意味を持つのであれば、時空を超えて2つの群が混合しているというのは不自然な前提である。そこで横軸に時系列的な意味を置くと、ある時点から状態が変化したものとして考えることができる。ここでその時点が「いつ」であるのかは不明であるが、わからないことを確率で表現するのが確率モデルのおもしろい点である。変化点を確率的パラメータとし、その前後で群が異なるというモデルは、変化点検出のモデリングと言われる。このモデリングはポリグラフ検査など、実践的な場面での利用価値も高い。

→Lee201709 の Pp.59–61, Matsuura201610 の Pp.238–245

折線回帰 平均点の位置が変わるだけでなく、変化の傾向が明らかな場合は線形モデルを当てはめることができる。変化点の前後で傾きが変わるような線形モデルは、折線回帰とも呼ばれる。折線回帰モデルの実装については、変化点と回帰モデルを組み合わせたあとで、折れる点を繋げる数学的補正を加えたモデルへと修正する。最後に、説明変数が時点であることから回帰分析の前提として標本の独立性が担保されていない問題を指摘する。

15.1.3 キーワード

- 混合分布モデル
- 変化点
- 折線回帰
- 時系列分析

15.2 授業情報

■コマの展開方法 講義/遠隔可/演習

15.2.1 予習・復習課題

■予習 混合分布モデルの応用になるので、混合分布モデルの基本的な書き方や解析方法について、第 13 講を復習しておくことが望ましい。

■復習 折線モデルが応用できそうなデータを見つけて、自分なりに実践してみると理解が深まるだろう。とくに折れる点が複数あるモデルや、折れる点の数を検出するモデルへと拡張するなど、モデル展開の可能性をかんがえることもできる。

16 確率的プログラミングの応用 3; 状態空間モデル

16.1 授業内容

16.1.1 科目の中でのこのコマの位置づけ

前時に時系列的なデータを導入したが、時系列的な性質を無視したモデリングになっていた。時系列的な分析手法は、心理学においてもウェアラブル端末の利用や SNS など公的なデータを分析することなどにも利用できるため、非常に有用なものになりうる。しかしデータの特徴として非独立性の問題、周期性やトレンドの存在などがあり、周波数解析をおこなったり多次元の行列分解などが必要である。中でも状態空間モデルは比較的シンプルであり、とくにベイズアンプローチで実装が容易になったと言えるだろう。

ここでは状態空間モデルの基本的な考え方を導入し、モデリングについて解説する。ここで状態と観測の分離を行い、とくに観測が行われていない点があっても分析できること、観測が行われていない点をパラメータとして保管することに言及する。観測が行われていない点が保管できるのであれば、未来の時点についても予測が可能になるということである。ホワイトノイズモデルでそれを行うと、確信区間が広がっていくことが観測される。そこでトレンドを入れたモデルにすることで、さらに予測の形を変えられることを学ぶ。

そのほかにも季節項など、時系列特有の情報を組み込むことや、二次元に展開することで空間データの分析にも応用できることに言及する。

16.1.2 コマ主題細目

時系列データの特徴 時系列的なデータがどのように得られ、どのようなシーンで利用可能であるかを概観する。ここで時系列データはサンプルの独立性が満たされていないという問題があるため単純な線形回帰は不適切であること、また周期性やトレンド、介入効果が出てくるまでの期間など独自に考えなければならぬことがいろいろ含まれている。これまで研究されてきた領域や研究方法について概観する。

状態空間モデル さまざまな分析方法がこれまで考えてこられているが、状態空間モデルはその中でも比較的簡単な数理的構造を持ち、またベイズアンモデリングを利用することでかつての分析モデルが必要としたスムージングなどを、特段意識することなく分析できる。状態と観測というモデルの基本構造を提示し、これらがどのように実装可能かをみる。

→Matsuura201610 の Pp.229-235, Baba201907 の第 5 部

欠測値の補間 観測時点には欠損が含まれることもあり、これをパラメータとして推定・補間することができる。またこれが可能であるということは、未来の時点が欠測値として考えれば予測ができることにもなる。プログラミングの工夫により、欠損を補間するようなコードの書き方を学ぶ。また単純なホワイトノイズモデルであればあまり予測として意味がないが、トレンドを考えることで時系列的な影響について考えることができる。

状態空間モデルの展開 状態空間モデルは、説明変数を加えた回帰モデルに応用したり、周期性をモデリングすることなども可能である。さらに時系列は一次元的であるが、二次元にも広げると空間分析にも利用が可能であることに言及する。これからの心理学は、時系列や空間など状況変数をより積極的に取り組んだモデルも利用するようになるだろう。

16.1.3 キーワード

- 状態空間モデル
- トренд
- 補間

16.2 授業情報

■コマの展開方法 講義/遠隔可/演習

16.2.1 予習・復習課題

■予習 時系列データを, 時間を独立変数とした回帰分析にすることでどういった問題があるのかについて, 回帰分析や確率モデルの前提などを考えて振り返っておくことが良い予習になるだろう。

■復習 身の回りの身近ところからでもデータを取ることができるのが, 時系列データのおもしろいところでもあるので, 応用可能なデータを探して分析してみると良い。可能であれば今日からでも, 時系列的なデータを取り始めると, 長期的に見て非常に興味深い分析ができるようになるだろう。