

Rendre utilisable une grande collection pédagogique

James Hanley

Department of Epidemiology, Biostatistics and Occupational Health

McGill University, Montréal, Québec, Canada

Présentation sur invitation

2 juin 2026

Congrès annuel de la Société statistique du Canada





Down the Victorian data mine

The Victorian data mine is a treasure trove of information, and it's being mined by a team of researchers at the University of Cambridge.

Kean's Victorian data mine is a treasure trove of information, and it's being mined by a team of researchers at the University of Cambridge. The data mine is a collection of Victorian-era data, including census records, birth and death records, and more. The data is being used to study the lives of Victorians and to understand the social and economic changes of the era.

The data mine is a collection of Victorian-era data, including census records, birth and death records, and more. The data is being used to study the lives of Victorians and to understand the social and economic changes of the era. The data mine is a treasure trove of information, and it's being mined by a team of researchers at the University of Cambridge.



Sondage anonyme de 30 secondes (facultatif)

Sondage anonyme de 30 secondes (facultatif)



PLAN

PLAN

- Brève bio en images [🔗* link](#)
- Les ressources en ligne issues de mon enseignement
[🔗* link](#)
- Vers un **catalogue** qui les rend repérables et utilisables par d'autres enseignants
 - Plan
 - Exemples
 - Appel à commentaires et à la collaboration





1



BSc (Mathematics/Statistics) 1968
MSc (Mathematics/Statistics) 1969

4



Elementary School: trip to Cork City ?1959

2

Secondary School, 1960-1965



St. Brendan's College, Killarney, Co. Kerry

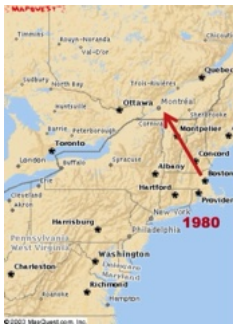
3



**PhD,
Statistics
1973**



1



4

**S.U.N.Y. / Buffalo
(Prof. Marvin Zelen)**

Eastern Cooperative
Oncology Group



RTOG
RADIATION THERAPY
ONCOLOGY GROUP

<http://ecog.dfci.harvard.edu>

<http://www.rtog.org>

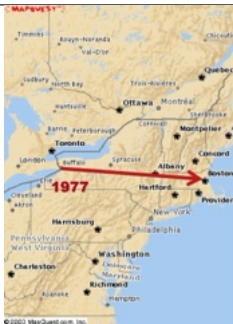
2



3



BIostatISTICS
Harvard School of Public Health
Boston, Massachusetts



Age 8
(1955)

GRADES 4-7
(another teacher)

GRADES 0-3
(1 teacher)

<https://jhanley.biostat.mcgill.ca/Reprints/ThingsTheyDontTeachYou-JHanley.pdf>



UNIVERSITY COLLEGE CORK
Coláiste na hOllscoile Corcaigh



BSc (Mathematics/Statistics) 1968
MSc (Mathematics/Statistics) 1969

<https://jhanley.biostat.mcgill.ca/Reprints/ThingsTheyDontTeachYou-JHanley.pdf>

MA PAGE D'ACCUEIL, 2026

MA PAGE D'ACCUEIL, 2026



James Hanley
Emeritus Professor
Biostatistics
McGill University



McGill
McGill University

Professor Emeritus [Epidemiology, Biostatistics and Occupational Health](#)

Webpage jhanley.biostat.mcgill.ca

E-mail james.hanley@mcgill.ca

Telephone +1 (514) 398-6270

c.v. / bio [longer](#) [short bio: pictures only](#)

Google Scholar [Google Scholar](#)

Search, via Google, for material on this site:

Welcome to this site, which I have maintained since the 1990s.

In July 2023 I moved it to a new server, with a new domain name.

Feel free to browse/use the materials. If you notice a link that is not working, or find any errors in the material (or if you want to discuss any of the topics) please feel free to contact me.

James Hanley
2023.07.19

Course Material

[[Back to Home Page](#)]

BIOS-601 [Epidemiology: Intro & Statistical Models \(Fall 2023\)](#)

EPID-405 [Critical Appraisal in Epidemiology \(Winter 2023\)](#)

BIOS-691 [Applied Statistics: hands-on data analysis \(Winter 2023 & 2021\)](#)

BIOS-624 [Data Analysis & Report Writing \(Fall 2016\)](#)

BIOS-602 [Epidemiology: Regression Models \(Winter 2015\)](#)

Workshop: [link](#)

2022: Computational and Data Systems Initiative

2017: Anatomy & Cell Biology

2015: Engineering (Concordia): Communication

2013: Bionano

Acad. 1/2 Day: [Regression/Multivariable Methods: Bi-3 Int.Med. 2011](#)

Acad. 1/2 Day: [Statistics in medical literature: Bi-3 Int.Med. 2012](#)

EPID-609 [Seminar: Epidemiology PhD students \(Fall 2021\)](#)

EPID-634 [Survival Analysis & Related topics Winter 2011](#)

EPID-613 [Introduction to Statistical Software Fall 2006](#)

EPID-694 [Statistical Inference II June 2004](#)

EPID-681 [Data Analysis II \(Winter 2004\)](#)

524-207 [Introduction to Epidemiology \(med2\) Fall 2002](#)

EPID-640 [Practices, Winter 2002](#)

MATH323 [Probability Theory May 2001](#)

EPID-607 [Principles of Inferential Statistics in Medicine Fall 2001](#)

EPID-610 [Lectures in GIS -- Nov 1999](#)

EPID-622 [Applications of Statistics](#)

EPID-626 [Risks and Hazards](#)

EPID-678 [Analysis of Multivariable Data](#)

MATH685 [Statistical Computing](#)

EPID-697 [Applied Linear Models](#)

[Back to Home Page](#)

UNE IDÉE DIRECTRICE

UNE IDÉE DIRECTRICE

**Commencer par la QUESTION,
pas par la méthode**

→ GRANDE COLLECTION EN LIGNE DE ...

→ GRANDE COLLECTION EN LIGNE DE ...

- Exercices développés localement (au fil des années)
- Rapports complets d'études réelles (historiques et récentes)
- Jeux de données associés
- Notes de cours (rédigées, pas seulement des diapositives)
- Articles pédagogiques (étudiants/collaborateurs + moi)
- Démonstrations en ligne de concepts (animations, vidéo)

→ GRANDE COLLECTION EN LIGNE DE ...

- Exercices développés localement (au fil des années)
- Rapports complets d'études réelles (historiques et récentes)
- Jeux de données associés
- Notes de cours (rédigées, pas seulement des diapositives)
- Articles pédagogiques (étudiants/collaborateurs + moi)
- Démonstrations en ligne de concepts (animations, vidéo)

NON CATALOGUÉ → invisible en pratique

Un ensemble d'exercices utilisés dans un cours [bios601/hw_sampling2023.pdf]

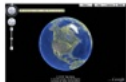
Un ensemble d'exercices utilisés dans un cours [bios601/hw_sampling2023.pdf]

Course BIOS601: **Warmup** ASSIGNMENT on Sampling and its Error – and Various Other Concepts. [underline = link]

Fall 2023 v. July 31

How Deep is the Ocean?

(A song – see Wikipedia – and an article – see [Significance](#) Dec, 2014)



1 What percentage of the world's surface is covered by water?

The data provided by the Scripps Institution of Oceanography [accessible via <https://jhanley.biotstat.ucgill.ca/bios601/surveys/Oceanography/>] can provide an answer, but some work is required on your part.

- i. In previous years, students drew a simple random sample of 200 locations on the Earth's surface,¹ and obtained from the `SRTM30 PLUS` database the land elevation or ocean depth at each of these. This year, to save some time², both the drawing of the sample, and the 200 database lookups have already been done for you – there are several `.csv` files (62 if JH has counted correctly) from previous years at the bottom of the 'Oceanography Data' webpage. To avoid picking the same datafile as another student, use your birthday `yyyymmdd` as the eight digit argument in the `set.seed()` function in R, and sample 1 from the numbers 1 to 62 using the `sample(1:62, 1)` call. From the 'readings' in your selected file, calculate a point estimate of the percentage.³ Also calculate a (probabilistic) margin of error (ME): do this by calculating a standard error, and multiplying it by say 1.96 so that you can make a probabilistic statement. [Again, use a *sensible no. of decimal places*]
- ii. Are you worried about the appropriateness of using 1.96 (and the Normal distribution) for the 95% confidence interval? Why/why not?

¹Ch 9 in Gelman & Nolan's *Teaching Statistics: A Bag of Tricks* (an ebook at McGill Library) has an interesting way of sampling, and other useful remarks on this problem. Today, one could simply zoom all the way out in Google Maps, and spin!

²You are still welcome to get your own 'from scratch.'

³Do not show off how many decimals you (R) can calculate. If your parents asked you what percentage you got, how many digits would you give them? Same applies to other Qs!

- iii. If you are – and even if you are not – find an online calculator/table that yields an 'exact' confidence interval. Compare the 'exact' interval with the 'approximate' one above.
- iv. Using what you are able to find online or from your textbooks, explain to a relative who is an engineer how exactly this 'exact' confidence interval is calculated and how the principles behind it differ from those behind the usual one. [We will come back to this in a later class].
- v. The root mean squared error includes both sampling variation and non-sampling errors. Your margin of error is limited to the sampling variation, and is modulated by the choice of 'n.' It does not include *non-sampling* errors.⁴ Describe one possible source of non-sampling error in this particular context of ocean depths.

Also, describe an unrelated example you would use to describe non-sampling errors to a lay person. [Internet searching is encouraged, but please cite the source if you found this example online, or in a textbook]

2 What is the average depth of the ocean?

- i. From the relevant observations (from among your 200), estimate the mean ocean depth, and calculate an accompanying ME.⁵ Even though there is a random component to it, pretend that the sample size was predetermined.
- ii. Are you worried about the appropriateness of using 1.96 (and the Normal distribution) for 95% confidence? Why/why not?

3 Ensuring that a sample of n' locations will yield $n = 200$ [or more] usable ones

- i. How big must n' be in order to have a good chance (say 80%) that it will yield at least 200 usable ones (i.e. ocean locations)?
- ii. What if you sampled sequentially until, at the n' -th draw, you reached the 200-th usable one? What distribution describes the random variable

⁴Some define a 'non-sampling' error as one that is not minimized by taking bigger and bigger n ; indeed, if there is some 'systematic' error in the measurements, taking an even bigger n will just make the answer more precisely wrong!

⁵In doing so, make sure to reduce your 'n' accordingly. Some students in previous years continued to use an n of 200!

Un ensemble d'exercices utilisés dans un cours [bios601/hw_sampling2023.pdf]

Course BIOS601: Warmup ASSIGNMENT on Sampling and its Error – and Various Other Concepts. [underline = link]

Fall 2023 v. July 31

n^* ? Calculate its 10-th and 90-th percentiles (pretend you know the value of the parameter that determines its distribution).^[6]

4 More efficient (or more practical) sampling strategies

(Very briefly) describe the circumstances^[7] in which a sampling scheme other than s.r.s (systematic, stratified, cluster) would offer either practical or statistical efficiency advantages; mention also the downsides of these schemes [textbook and internet searching encouraged – if you acknowledge the source!].

5 Oh Oh

(a) One way to obtain random (λ =longitude, ϕ =latitude) locations is as $\lambda \sim U(-\pi, \pi)$, and $\phi \sim \text{pdf}(\phi) = (1/2) \cos(\phi)$ on $(-\pi/2, \pi/2)$.

Figure 4B (p. 10) was considered too technical for the Significance article. It is included on p.10 below to help explain how one could sample the ϕ 's. Think of a longitudinal-based section of a (perfectly spherical!) orange, and of how wide it is at 'latitude' ϕ compared with how wide it is at 'latitude' 0 (on the equator). It is the same as the relationship between the west-east distance between two locations at the same latitude (e.g., $\phi = 45.5^\circ$), but say 1 degree longitude apart, and the (approx. 100km) west-east distance between two locations on the equator, also 1 degree longitude apart. If we take the distance at the equator to be 1, the distance at 'latitude' ϕ is $\cos(\phi)$. Thus, in any chosen longitude-based section the number of sampled locations at latitude ϕ should be $\cos(\phi)$ times the number sampled at the equator.

- Use your dataset of 200 to check that the random locations produced by this method [implemented in the R code used to create the personalized datasets for question 1(i)] appear to be sensible.

[Question 10, *The Locations of the Stars*, takes up the question of randomness, from a different viewpoint!]

(b) A researcher spent the entire research budget on a sample of 200 locations, using $\lambda \sim U(-\pi, \pi)$, but $\phi \sim U(-\pi/2, \pi/2)$.

- Explain why this sampling scheme is flawed. [Gelman and Nolan have a few words on this]. Are the resulting data worthless? Or, do you think we could recover something from them?

- Using the information in (a), suggest a way to correct for the researcher's oversampling of locations further from the equator.

(c) Search online (or in your textbooks) for ways to draw random samples from a non-uniform continuous distribution. List ones that are easy to implement when only the (i) pdf, (ii) CDF has a closed form.

(d) The rationale behind the 'inverse CDF' method is often missed – and not easily recalled years later – if students go through the 'proof' as a mere 'math-stat' or calculus exercise.

Figure 4B (p.10) tries to explain the 'inverse CDF' method in pictures rather than via calculus.^[8]

Pages 11–12 are notes from 2010, with yet another plot of *west-east lines laid end to end* – another attempt to 'explain' it in this same 'sampling latitudes' context.

The attempt on page 13 uses an unnamed continuous random variable, but starts with a simple discrete version that might make the methods more intuitive.

• Now the test of whether any of these three attempts succeeded: **in your own words**, explain to that same relative of yours how exactly the inverse CDF methods work. If you don't like the examples/explanations JH has provided, feel free to make up your own.^[9]

^[6]This article https://sanjay.khastgir.org/teaching/bios601/hw_sampling2023.pdf originally had the data-mining challenge, and a description of the method to generate random locations. But the diagram (now on page 7 below) was considered too complex and too technical for the Significance Magazine readership.

^[7]In the past, JH has heard a teacher start by asking students whether in a distribution – any distribution – there are more/fewer people between the 55th and 56th percentile than there are between the 5th and 6th, or 95th and 96th? This teacher was also quite fussy about words, and about using the word 'percentile' correctly, so he would probably have taken exception to JH's saying *percentiles* are numbered 1–100.

^[8]R has 'exact' d- p- and q- functions for this distribution, but – given the numbers involved – the calculations can also be reasonably approximated if you know just its mean and variance.

^[9]The Cross-Canada Survey of Radon Concentrations in Homes [Resources] might help.

Question 1 : Quel % de la surface de la Terre est recouvert d'eau?

Question 1 : Quel % de la surface de la Terre est recouvert d'eau?

- i. In previous years, students drew a simple random sample of 200 locations on the Earth's surface,¹ and obtained from the SRTM30 PLUS database 933 million locations the land elevation or ocean depth at each of these. This year, to save some time², both the drawing of the sample, and the 200 database lookups have already been done for you – there are several .csv files (62 if JH has counted correctly) from previous years at the bottom of the 'Oceanography Data' webpage. To avoid picking the same datafile as another student, use your birthday `yyyymmdd` as the eight digit argument in the `set.seed()` function in R, and sample 1 from the numbers 1 to 62 using the `sample(1:62,1)` call. From the 'readings' in your selected file, calculate a point estimate of the percentage.³ Also calculate a (probabilistic) margin of error (ME): do this by calculating a standard error, and multiplying it by say 1.96 so that you can make a probabilistic statement. [*Again, use a sensible no. of decimal places*]
- ii. Are you worried about the appropriateness of using 1.96 (and the Normal distribution) for the 95% confidence interval? Why/why not?

¹Ch 9 in Gelman & Nolan's *Teaching Statistics: A Bag of Tricks* (an ebook at McGill Library) has an interesting way of sampling, and other useful remarks on this problem. Today, one could simply zoom all the way out in Google Maps, and spin!

²You are still welcome to get your own 'from scratch.'

³Do not show off how many decimals you (R) can calculate. If your parents asked you what percentage you got, how many digits would you give them? Same applies to other Qs!

PARTIR DE LA FAÇON DONT LES
ENSEIGNANTS EFFECTUENT LEURS
RECHERCHES

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**
 - Concept (quelle notion statistique?)

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**
 - Concept (quelle notion statistique?)
 - Niveau (pour quels étudiants?)

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**
 - Concept (quelle notion statistique?)
 - Niveau (pour quels étudiants?)
 - Type (comment l'utiliser?)

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**
 - Concept (quelle notion statistique?)
 - Niveau (pour quels étudiants?)
 - Type (comment l'utiliser?)
 - Contexte (de quoi s'agit-il?)

PARTIR DE LA FAÇON DONT LES ENSEIGNANTS EFFECTUENT LEURS RECHERCHES

- **Les enseignants travaillent par besoin :**
 - Concept (quelle notion statistique?)
 - Niveau (pour quels étudiants?)
 - Type (comment l'utiliser?)
 - Contexte (de quoi s'agit-il?)
 - Accroche (qu'est-ce qui capte l'attention?)

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

- **Concept:** échantillonnage, variabilité, estimation

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

- **Concept:** échantillonnage, variabilité, estimation
- **Niveau:** introductif / intermédiaire

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

- **Concept:** échantillonnage, variabilité, estimation
- **Niveau:** introductif / intermédiaire
- **Type:** discussion en classe, devoirs

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

- **Concept:** échantillonnage, variabilité, estimation
- **Niveau:** introductif / intermédiaire
- **Type:** discussion en classe, devoirs
- **Contexte:** couverture des océans, données massives

UN EXEMPLE – PLUSIEURS POINTS D'ENTRÉE

- **Concept:** échantillonnage, variabilité, estimation
- **Niveau:** introductif / intermédiaire
- **Type:** discussion en classe, devoirs
- **Contexte:** couverture des océans, données massives
- **Accroche:** ordre de grandeur surprenant, données réelles

VERS UN CATALOGUE CONVIVIAL : UNE STRUCTURE SIMPLE

VERS UN CATALOGUE CONVIVIAL : UNE STRUCTURE SIMPLE

- Concept
- Niveau
- Type
- Contexte
- Accroche

VERS UN CATALOGUE CONVIVIAL : UNE STRUCTURE SIMPLE

- Concept
- Niveau
- Type
- Contexte
- Accroche

Chaque facette a un petit vocabulaire contrôlé

VERS UN CATALOGUE CONVIVIAL : UNE STRUCTURE SIMPLE

- Concept
- Niveau
- Type
- Contexte
- Accroche

Chaque facette a un petit vocabulaire contrôlé

- Ce n'est pas parfait
- Ce n'est pas exhaustif
- **Mais c'est cohérent**

DE LA RECHERCHE À L'UTILISATION

Recherche → Parcourir → Décider → Utiliser

Catégories de **CONTEXTE** proposées

Catégories de **CONTEXTE** proposées

1. Santé / médecine
2. Biologie
3. Environnement
4. Éducation
5. Démographie / population
6. Technologie / ingénierie
7. Social
8. Mathématiques

Catégories de **TYPE** proposées

Catégories de **TYPE** proposées

1. Question/discussion en classe
2. Exercice
3. Projet / tâche prolongée
4. Exploration de données
5. Illustration de concept
6. Évaluation (quiz/examen)
7. Étude de cas

Catégories de **NIVEAU** proposées

Catégories de **NIVEAU** proposées

1. Introductif
2. Intermédiaire
3. Avancé

Catégories d'**ACCROCHE** proposées

Catégories d'**ACCROCHE** proposées

- Résultat surprenant
- Résultat contre-intuitif
- Cas historique réel
- Données massives
- Pertinence politique
- Récit humain
- Défi de mesure
- Dispositif / démo / application intéressante

Catégories de **CONCEPT** - base de départ

Catégories de **CONCEPT** - base de départ

- Données et mesures
 - types de variables | erreur de mesure | conséquences
- Statistiques descriptives
 - quantiles | emplacements | diffusion | forme | mauvaises interprétations
- Plan d'étude et échantillonnage
 - expérimental ou par observation | méthodes d'enquête
- Probabilité et distributions
 - calculs | modèles | discret | continu
- Inférence statistique
 - estimation | intervalles de confiance | tests d'hypothèse
- Modélisation et relations
 - corrélation | régression | hypothèses

CE QUE LE SYSTÈME ESSAIE DE FAIRE

CE QUE LE SYSTÈME ESSAIE DE FAIRE

- Non pas : trouver le meilleur élément
- Mais : réduire le choix à un petit ensemble exploitable

Objectif: de 5 à 10 éléments pertinents

CE QUE LE SYSTÈME ESSAIE DE FAIRE

- Non pas : trouver le meilleur élément
- Mais : réduire le choix à un petit ensemble exploitable

Objectif: de 5 à 10 éléments pertinents

Pour que :

- leurs descriptions puissent être parcourues rapidement
- l'enseignant puisse faire un choix final

COMMENT UN ENSEIGNANT EFFECTUERAIT UNE RECHERCHE

COMMENT UN ENSEIGNANT EFFECTUERAIT UNE RECHERCHE

Concept		Level		Type
Data & Measurement	☒	Introductory		Short discussion/prompt
Descriptive Statistics		Intermediate	☒	Exercise
☒ Study Design & Sampling		Advanced		Project / extended task
Probability & Distributions				Data exploration
Statistical Inference				Concept illustration
Modeling & Relationships				Assessment item (quiz/exam)
				Case study
Context		Hook		
Health / Medicine		Surprising result		
Biology		Counterintuitive finding		
Environment		Real historical case		
Education	☒	Big data		
Demography / Population		Policy relevance		
Technology / Engineering		Human story		
Social		Measurement challenge		
Mathematical		Neat device/app/demo/story		

RÉSULTATS RETOURNÉS

RÉSULTATS RETOURNÉS

- What percentage of the Earth's surface is covered by water?
- What is the average depth of the ocean?
- What happened to this person's physical activity over the years 2010-2021?
- How does the fuel economy of the newer family van compare with its predecessor?
- When will the ice break up this year? How close were the 0.25M bets?
- Do children provided with milk at school grow more?
- Does the switch to/from daylight savings time affect rates of motor vehicle accidents?
- Mystery Canadian Data: what was going on?
- How to determine the efficacy of the first iteration of the Salk polio vaccine?

CELUI-CI SEMBLE PROMETTEUR...

- What percentage of the Earth's surface is covered by water?
- What is the average depth of the ocean? 
- What happened to this person's physical activity over the years 2010-2021?
- How does the fuel economy of the newer family van compare with its predecessor?
- When will the ice break up this year? How close were the 0.25M bets?
- Do children provided with milk at school grow more?
- Does the switch to/from daylight savings time affect rates of motor vehicle accidents?
- Mystery Canadian Data: what was going on?
- How to determine the efficacy of the first iteration of the Salk polio vaccine?

DÉTAIL D'UNE FICHE DU CATALOGUE

DÉTAIL D'UNE FICHE DU CATALOGUE

Titre: What is the average depth of the ocean?

Concepts: Study Design & Sampling; Statistical Inference

Niveau: Introductory

Type: Exercise; Case study

Contexte: Environment

Accroche: Big data

Matériel: Background articles | exercise | samples | teacher notes | code

Math / stats avancées No

Durée: 15-30 minutes

Accroche spécifique: 19th century scientist & model; 21st century validation($N = 9.3 \times 10^8$)

URL(s) etc. : Click

DÉTAIL D'UNE FICHE DU CATALOGUE

Titre: What is the average depth of the ocean?

Concepts: Study Design & Sampling; Statistical Inference

Niveau: Introductory

Type: Exercise; Case study

Contexte: Environment

Accroche: Big data

Matériel: Background articles | exercise | samples | teacher notes | code

Math / stats avancées No

Durée: 15-30 minutes

Accroche spécifique: 19th century scientist & model; 21st century validation($N = 9.3 \times 10^8$)

URL(s) etc. : Click

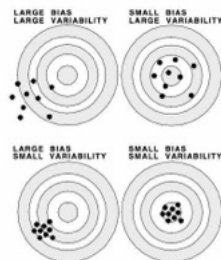
Révélation progressive: parcourir → cliquer → adapter → utiliser

2

MESURE

Course BIOS601: The Quality of Measurements and the Effects of Measurement Error.

Fall 2023, v. 09.06



Being approximately correct and being precisely wrong

PREFACE

Few general statistics textbooks¹ deal with the topic of measurement error. However, it is an important topic. Measurement error is present in all scientific work; while it can sometimes be lessened with careful planning and extra effort, it cannot be completely avoided. In some contexts it just adds noise and makes signals harder to detect. *But in others, its effects are both more subtle and more serious:* it systematically shifts the parameter estimates produced by models.

Its effects become more complicated and unpredictable in larger statistical models with two or more measured-with-error variables. We will restrict attention here to simple comparative situations involving one X and one Y . We

Exercises 1-8 (sur 30)

Exercises 1-8 (sur 30)

EXERCISES

1. Refer to the descriptions of the SMOG index, the Fry method, the Flesch Reading Ease, and the Flesch-Kincaid Grade Level, for measuring readability (under Resources for Measurement/Surveys).^[10] [Also, see Q29]

For the article or text you have chosen (as per discussion in class), randomly select three separate 100 word passages, and use this *set of three passages* to measure the readability (F_1) using the Fry graph. Rather than do so manually, you can use the SMOG calculator to determine the average number of sentences and syllables per hundred words. Repeat the readability measurement (F_2) with a second *different* set of three passages. Repeat once more (F_3), using a *third set*.

Using these same three sets, calculate the SMOG index, the Flesch Reading Ease, and the Flesch-Kincaid Grade Level.

For each index, use the 3 estimates to calculate the standard error of measurement, and the coefficient of variation. Comment.

2. Propose a method to assess the validity of a readability index.
3. [m-s] Derive the link between the standard error of measurement and the (intraclass correlation) reliability coefficient [last line, column 1, p16 in the notes above]. *Hint: it's simply a matter of using the definition of R.*
4. [m-s] Exercise in section 3 (p. 23 of notes) of Relationship between test-retest correlation and ICC(X) [In notes on Effect of Errors in X and Y on measured correlation and slope]
5. [m-s] Exercise section 4: Relationship between correlation(X, X') and ICC(X) [ibid.]
6. Francis Galton (1822-1911) found that the correlation between (*self-reported*) parental and (adult) offspring heights was strongest for the one between father and son [0.396 ± 0.024], and weakest for the one between mother and daughter [0.284 ± 0.028]. Given the way he obtained the measurements, can you imagine why this was? ^[11] [It was 0.302 ± 0.027 for mother & son; 0.360 ± 0.026 for father & daughter.]

¹⁰ToolCheck (<https://teachmean.com/2010/07/20/toolcheck/>) "sounded" like an interesting tool; it's not clear if "make it" commercially, or was bought by another company!

¹¹After you have thought about it for a while, and looked carefully at Galton's Notebook, you might wish to compare your answer with Karl Pearson's explanation: "Why Galton got different parent-offspring correlations in heights" <https://janley.biostat.mcgill.ca/bios601/Surveys/pearson1930v13e1ch4p17-18.pdf> and also look at "why he (KP) got larger

7. *Bridging the physical- and the psycho-metric:* The notes on "Increasing Reliability by averaging several measurements" on the right hand column of page 13 of JH's notes on Quantifying Reliability give the formula for the so-called "Stepped-Up Reliability". In psychometrics (where the number of items on a test serves as the "several measurements") this formula serves as the basis for the "Spearman-Brown prediction formula".^[12]

[m-s] Invert the formula on p.13 to derive the one on the right hand column of p.10 for the Spearman-Brown prediction formula relating the reliability of two versions of a test, one with N times more items than the other.

8. You are trying to estimate, from imperfect observations of F and C , the values of the two coefficients B_0 and B_1 in the temperature relation $F = B_0 + B_1 \times C$.

For each of the following situations, and using the true values $B_0 = 32$ and $B_1 = 9/5 = 1.8$, simulate^[13] 1000 datasets and investigate the behaviour of the 1000 estimates, b_0 and b_1 , of B_0 and B_1 . In each simulation, use samples of size $n = 4$, with temperatures of $C = 14, 16, 18$ and 20 .

- (a) C measured perfectly, F measured with $\epsilon_F \sim \text{Gaussian}(\mu = 0, \sigma_{\epsilon_F} = 1)$ errors that are independent of F . Check – formally, using a test (or CI) based on the mean of the 1000 estimates – for evidence of bias in b_1 . Also check whether the empirical variance of b_1 agrees with that given by the theoretical formula, namely

$$\text{Var}(b_1) = \sigma_{\epsilon_F}^2 / \sum (x - \bar{x})^2.$$

- (b) F measured perfectly, C measured with $\epsilon_C \sim \text{Gaussian}(\mu = 0, \sigma_{\epsilon_C} = 1)$ errors that are independent of C [Classical type error: someone else chose situations when C was indeed exactly 14, 16, etc, but didn't tell you what C was, and instead asked you to independently record C using your own imperfect instrument, and to use your recordings of C in your estimation of the equation]. Again, formally test for evidence of bias in b_1 .

ones" <https://janley.biostat.mcgill.ca/bios601/Surveys/pearson1930p377-378.pdf> using this protocol (p358-) <https://janley.biostat.mcgill.ca/bios601/Surveys/Pearson1930.pdf> is the "Measurement - Lecture Notes, etc" section of the bios601 resources page for Measurement.

¹²Wikipedia has an entry called "Spearman-Brown prediction formula".
¹³If new to simulations, see "Computer code to simulate datasets with measurement error" <https://janley.biostat.mcgill.ca/bios601/Surveys/FandC.k.txt> at the bottom of the Resources webpage for measurement/surveys. It gives some 'starter' computer code, which you can modify to suit.

Corrélations dans les données issues du crowdsourcing (en échange de > 100,000 \$)

Corrélations dans les données issues du crowdsourcing (en échange de > 100,000 \$)

6. Francis Galton (1822-1911) found that the correlation between (*self-reported*) parental and (adult) offspring heights was strongest for the one between father and son [0.396 ± 0.024], and weakest for the one between mother and daughter [0.284 ± 0.028]. Given the way he obtained the measurements, can you imagine why this was? ¹¹ [It was 0.302 ± 0.027 for mother & son; 0.360 ± 0.026 for father & daughter.]

¹¹After you have thought about it for a while, and looked carefully at Galton's Notebook, you might wish to compare your answer with Karl Pearson's explanation: "Why Galton got different parent-offspring correlations in heights" <https://jhanley.biostat.mcgill.ca/bios601/Surveys/pearson1930vol13ach14p17-18.pdf> and also look at "why he (KP) got larger ones" <https://jhanley.biostat.mcgill.ca/bios601/Surveys/PearsonBka1902pp377-378.pdf> using this protocol (p358-) <https://jhanley.biostat.mcgill.ca/bios601/Surveys/PearsonLee1903.pdf> in the 'Measurement - Lecture Notes, etc' section of the bios601 resources page for Measurement.

¹²Wikipedia has an entry called 'Spearman Brown prediction formula'.

¹³If new to simulations, see "Computer code to simulate datasets with measurement error" <https://jhanley.biostat.mcgill.ca/bios601/Surveys/FandC.R.txt> at the bottom of the Resources webpage for measurement/surveys. It gives some 'starter' computer code, which you can modify to suit.

	Fils	Fille
Père	0.396 ± 0.024	0.360 ± 0.026
0.10		
Mère	0.302 ± 0.027	0.284 ± 0.028

FICHE DE CATALOGUE

Title: Height: why is father-son correlation larger than mother-daughter?

Concept: Data & Measurement; Modeling & Relationships

Level: Introductory; Intermediate

Type: Short discussion/prompt; Concept illustration; Exercise

Context: Biology

Hook: Surprising result; Measurement challenge

Materials: Background articles | data | exercise | teacher notes | code

Calculus/Math-stat Depends

Length: 15-30 minutes

Specific hook: Attenuation: caused by self-reported rather than measured heights
Article: teaching multiple regression with family data

URL(s) etc. : [Click](#)

3

Exemple 3 - une nouvelle tournure d'un classique intemporel

Exemple 3 - une nouvelle tournure d'un classique intemporel

2012

Material for Academic 1/2 day:
Statistics in the Medical Literature and in Medical Practice
October 18, 2012

'Averages'

- [Why do statisticians commonly limit their inquiries to averages?](#)
- [Salaries of professional baseball players 1994 Raw Data\(.xls\)](#)
- ["Average" height of basketball team](#)
- Number of authors on medical articles (Fletcher NEJM 1979)

Year	No. articles examined	No. au's mean	SD	No. Subjects Median
1946	151	2.0	1.4	25
1956	149	2.3	1.6	36
1966	157	2.8	1.2	16
1976	155	4.9	7.3	30

- [A new sex survey of American men \(1993\)](#)
- [The Myth, the Math, the Sex - NY Times](#)
- [The Median, the Math and the Sex - NY Times](#)
- [Statistical methods Radiology readers need to understand](#)
- [Visualizing the Median as the Minimum-Deviation Location Fig 1 \[colour\]; java applet](#)

There are Probabilities and Probabilities

- [Different Probabilities](#)

The Place of Statistical Methods in Radiology: & the Bigger Picture (Int. Med!)

- [article](#)

'P-values'

- [What the P-Value IS and IS NOT /](#)
- [L-Phenylalanine Mustard \(L-PAM\) in Management of Primary Breast Cancer \(Fisher's Exact Test\)](#)
- [The truth about doctors' handwriting: a prospective study](#)

[Introduction](#) [Individual Patient](#) [Imprecision](#) [CI's](#) [P-Values etc.](#) [Applications](#) [Summary](#)

P-Values and Statistical 'Tests'

"P-Value"

Defⁿ. A probability concerning the observed data, calculated under a **Null Hypothesis** assumption, i.e., assuming that the only factor operating is sampling or measurement variation.

Use To assess the evidence provided by the sample data in relation to a pre-specified claim or 'hypothesis' concerning some parameter(s) or data-generating process.

Basis As with a confidence interval, it makes use of the concept of a *distribution*.



45/61

Example 3 - a new twist on an old classic

Introduction

Individual Patient

(Im)precision

CI's

P-Values etc.

Applications

Summary

Example 1 – from *Design of Experiments*, by R.A. Fisher

Lady claims she can tell which was poured first...



BLIND TEST

Lady Says		4	0	4	4
		0	4	4	4
		4	4		

“Null Hypothesis” (H_{null}): she can not tell them apart.

Blind test is equivalent to being asked to say **which 4** of the following 8 Gaelic words are the **correctly spelled** ones. You are told that **4 are correctly spelled & 4 are not**.

1	2	3	4	5	6	7	8
madra	olscoil	cathiar	tanga	doras	cluicha	féar	bóthar

“Alternative” Hypothesis (H_{alt}): she can (can you think of another “H” ?).

Example 3 - a new twist on an old classic

Introduction

Individual Patient

(Im)precision

CI's

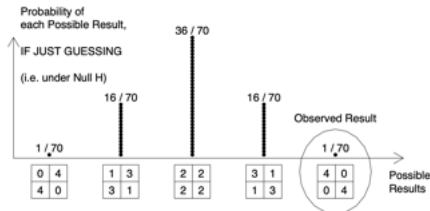
P-Values etc.

Applications

Summary

The evidence provided by the test

- Rank possible test results by degree of evidence against H_{null} .
- "P-value" is the probability, calculated under null hypothesis, of observing a result as extreme as, or more extreme than, the one that was obtained/observed.



In this e.g., observed result is the most extreme, so

$$P_{value} = \text{Prob}[\text{correctly identifying all 4, IF merely guessing}] = 1/70 = 0.014.$$

- Interpretation of such data often rather simplistic, as if these *data alone* should *decide*: i.e. if $P_{value} < 0.05$, we 'reject' H_{null} ; if $P_{value} > 0.05$, we don't (or worse, we 'accept' H_{null}). Avoid such simplistic 'conclusions'.

FICHE DE CATALOGUE

FICHE DE CATALOGUE

Title: A new twist on Fisher's 'lady tasting tea'

Concept: Probability & Distributions; Statistical Inference

Level: Introductory

Type: Concept illustration; Short discussion/prompt

Context: Biology

Hook: Neat device/app/demo/story

Materials: Course Notes;

Calculus/Math-stat No

Length: 5-10 min

Specific hook: Elegant, adaptable intro to idea of null hypothesis

URL(s) etc. : [Click](#)

O'SHEAS™

Rolling in Clover

20-Spot



MARK 20 NUMBERS
18 DIFFERENT WAYS TO WIN
\$3.00 Minimum Bet

TICKET MAY BE PLAYED FOR
ANY MULTIPLE OF \$3.00

CATCH	YOU WIN
-------	---------

0 of 20	300.00
1 of 20	6.00
2 of 20	FREE PLAY
3 of 20	FREE PLAY
4 of 20	-0-
5 of 20	-0-
6 of 20	-0-
7 of 20	FREE PLAY
8 of 20	6.00
9 of 20	15.00
10 of 20	30.00
11 of 20	120.00
12 of 20	600.00
13 of 20	3000.00
14 of 20	7000.00
15 of 20	13,000.00
16 of 20	19,000.00
17 of 20	30,000.00
18 of 20	40,000.00
19 of 20	50,000.00
20 of 20	50,000.00

4

Exemple 4: 1973

Example 4: 1973

THE TEACHER'S CORNER

Transformations Made Transparently Easy, or, So That's What a Jacobian Is!

DONALD G. WATTS*

*Dept. of Mathematics, Queen's Univ., Jeffery Hall, Kingston, Ont., Canada. The author has copyrighted this article in Canada.

The American Statistician, February 1973, Vol. 27, No. 1

Summary Simple teaching aids are used to help students understand and appreciate what a transformation is and what a Jacobian is. The teaching aids are effective and easy to make.



1. Introduction

One of the topics in introductory courses in probability and statistics which is often badly received is "transformations." Often times the student is not convinced of what he has been told, because no real transformations are shown, and those which *are* shown are in one dimension. The extension to more than one dimension is often done, under the guise of generality, with crude and sloppy amoeba-like sketches on the board. This is unsatisfactory because the student can not relate the sketch to his experience, and because the board is only two-dimensional. Where is the probability density function, and how is *it* being affected?

This note describes three highly useful devices for teaching transformations, all centered, quite naturally and obviously, around an overhead transparency projector. We are not concerned, here, about the *reasons* for transformation, only the *effects* of transformation.

Example 4: 1973



a) Showing $y = x^2$
and $f_X(x)$



b) Tipping
from x into y



c) Tipping into
 y from x
(reversed)



d) Showing $x = (y)^{1/2}$
and $f_Y(y)$

For the continuous case, we also need to produce a sample space and probability. We can, of course, only approximate the continuous case. The approach I have used is to make a sandwich of $\frac{1}{8}$ -inch plastic sheets separated by $\frac{1}{8}$ -inch plastic rods, as shown in Plate 1. The probability is fine sand.** The transformation shown is $Y = X^2$.

** The author is grateful to the University of Wisconsin Statistics Department for leaving one of their sand-filled ashtrays unguarded.

PLATE 1: THE TRANSFORMATION SANDWICH.

Exemple 4: 2006

Example 4: 2006

The PDF of a Function of a Random Variable: Teaching its Structure by Transforming Formalism into Intuition

James A. HANLEY and Dana TELTSCHE

The American Statistician, February 2006, Vol. 60, No. 1

In many instances, the probability density function (pdf) of a function of a random variable is obtained from the pdf of the random variable, the inverse function and the derivative of this inverse. The formula tends to be memorized rather than fully understood. This article describes how we teach the structure of the new pdf by (a) treating the problem as a change in scale, rather than a "change in variable"; (b) appealing to the concept of "conservation of probabilities"; (c) using the physical analogy of "pouring" probability mass from one set of "containers" to another; and (d) dealing fully with linear transformations before considering nonlinear ones.

KEY WORDS: Graphics; Physical representation; Probability mass; Scales; Transformations.

1. INTRODUCTION

The increasing availability of digital tools is an opportunity for teachers and authors to re-examine how statistical concepts are best presented. We consider the topic found under the rubric "transformation" or "change of variable." The goal is to understand the form of the probability density function (pdf) of a function h (taken, for now, to be monotonic) of a continuous random variable. The way the topic is presented today suggests that no matter whether teachers or authors favor the traditional

ously, around an overhead transparency projector." Sadly, none of the textbooks we examined have taken up his ideas, and the topic of transformations is often just as "badly received" today as it was when Watts wrote.

We describe ways to make the structure of the pdf formula more intuitive. A commonly used example is given in Section 3, and a practical example in Section 4, but we begin with two simpler ones to show the concepts involved. We deal with the bivariate situation in Section 5.

2. SIMPLE, LOCAL EXAMPLES WITH LINEAR TRANSFORMATIONS

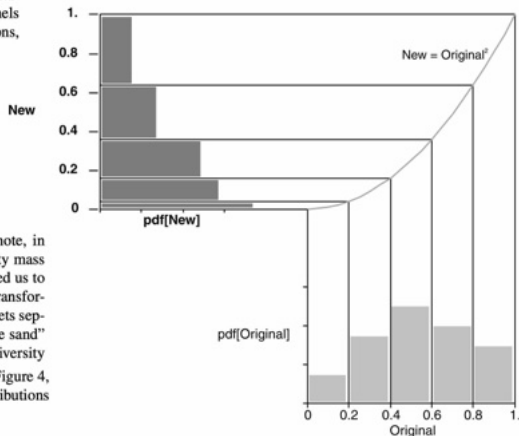
Suppose our focus is on day-to-day variations in temperature (T) in Montreal during a certain time of the year. Even though the early temperatures in the series were originally measured in degrees F, Canada switched to the metric system in 1975, and now the summaries of the entire series from the past 100 years are reported in degrees C. For the sake of illustration, suppose that on this Celsius scale, the distribution is well approximated by $T_C \sim N[\mu = 5, \sigma = 4]$. The pdf of this random variable is shown on the left panel of Figure 1, using the temperature scale $-5C$ to $+15C$, shown in plain typeface, and with the corresponding pdf scale shown on the left vertical axis, running from 0 to 0.1.

Those not familiar with the Celsius scale will wish to change the temperatures to Fahrenheit, using the conversion

$$T_F = 32 + (9/5)T_C.$$

Example 4: 2006

The upper panel of Figure 3 shows that the same reasoning applies to situations where the range of the new variable is different from that of the original. One can also use the pair of panels to show that one can carry out a sequence of transformations, just as children do when playing with water.



After we had completed an earlier version of this note, in which we had used the imagery of “pouring” probability mass from one set of containers to another, a colleague pointed us to the article by Watts (1973). His device for univariate transformations consisted of “a sandwich of 1/8-inch plastic sheets separated by 1/8-inch plastic rods. The probability was fine sand” (taken from an unguarded sand-filled ashtray in the University of Wisconsin Statistics Department). As illustrated in Figure 4, by tipping his device, he could “pour” probability distributions from one scale to another.

Figure 4. “Pouring probability,” after Watts (1973). Arbitrary original distribution, and transformation $New = h[Original] = Original^2$. By tipping the device, the probability mass in the original distribution (lower right) is poured into the new probability distribution (on its side, upper left).

Example 4: 2026

Example 4: 2026

`https://www.youtube.com/watch?v=VPQvRZ_jkQ0`

chaîne : `https://www.youtube.com/@YouStatistics`

Example 4: 2026

`https://www.youtube.com/watch?v=VPQvRZ\_jkQ0`

chaîne : `https://www.youtube.com/@YouStatistics`



FICHE DE CATALOGUE

FICHE DE CATALOGUE

Title: Pouring probability mass: the pdf of a function of a random variable

Concept: Probability & Distributions ; Modeling & Relationships

Level: Introductory

Type: Concept illustration; Exercise

Context: Mathematical

Hook: Neat device/app/demo/story

Materials: Course Notes (Math323); 1973 and 2006 teaching articles; animation
Calculus/Math-stat Yes
Length: 5-10 min. (demo+discussion); 30 min. (summarize 1973/2006 articles)

Specific hook: visual transformation; physical analogy; animation

URL(s) etc. : Click

Alors ... ?

OÙ EN SOMMES-NOUS

OÙ EN SOMMES-NOUS

- Une structure fonctionnelle pour décrire chaque élément
- Une façon de réduire une grande collection à **5-10 options exploitables**
- Un prototype de recherche et d'utilisation

OÙ EN SOMMES-NOUS

- Une structure fonctionnelle pour décrire chaque élément
- Une façon de réduire une grande collection à **5-10 options exploitables**
- Un prototype de recherche et d'utilisation

Mais plusieurs choix de conception restent encore à faire :

OÙ EN SOMMES-NOUS

- Une structure fonctionnelle pour décrire chaque élément
- Une façon de réduire une grande collection à **5-10 options exploitables**
- Un prototype de recherche et d'utilisation

Mais plusieurs choix de conception restent encore à faire :

- Jusqu'où standardiser le vocabulaire?
- Quels champs doivent être fixes vs flexibles?
- Qu'est-ce qui est essentiel vs agréable à avoir ?

OÙ EN SOMMES-NOUS

- Une structure fonctionnelle pour décrire chaque élément
- Une façon de réduire une grande collection à **5-10 options exploitables**
- Un prototype de recherche et d'utilisation

Mais plusieurs choix de conception restent encore à faire :

- Jusqu'où standardiser le vocabulaire?
- Quels champs doivent être fixes vs flexibles?
- Qu'est-ce qui est essentiel vs agréable à avoir ?

Si, ensemble, nous pouvions l'affiner sur une collection, nous pourrions l'étendre à d'autres.

Si, ensemble, nous pouvions l'affiner sur une collection, nous pourrions l'étendre aux autres

INVITATION

INVITATION

- Testez cette structure sur votre propre matériel
- Proposez de meilleurs termes
- Aidez à préciser l'essentiel

INVITATION

- Testez cette structure sur votre propre matériel
- Proposez de meilleurs termes
- Aidez à préciser l'essentiel

Une bonne couche de catalogue :



de nombreux supports pédagogiques
de qualité deviennent accessibles et
utilisables

Le MEILLEUR testeur de charge : pas moi, VOUS!

Le MEILLEUR testeur de charge : pas moi, VOUS!

J'aimerais trouver du matériel pour...

Le MEILLEUR testeur de charge : pas moi, VOUS!

J'aimerais trouver du matériel pour...



Vos réponses aideront à vérifier si cette structure correspond à vos besoins d'enseignement

MERCI

DES QUESTIONS?

DES COMMENTAIRES?

DES SUGGESTIONS?