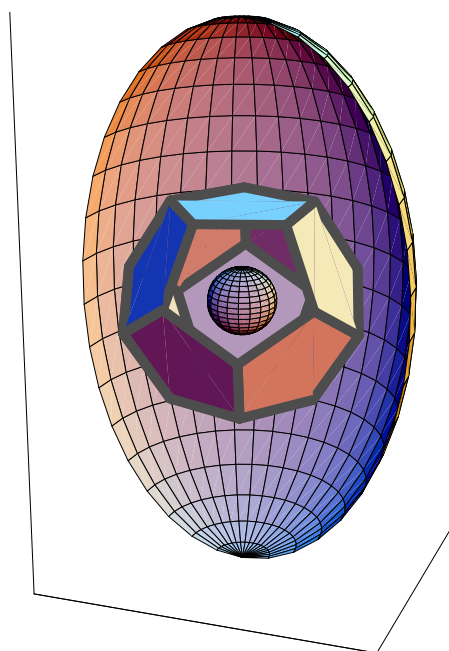


Script ◇ Math ◇ Ing
◇ Wahrscheinlichkeit & Statistik ◇
◇ Probabilité & Statistique ◇
kurz & bündig ◇ concis



Scripta bilingua

von • *de*

Rolf Wirz

Berner Fachhochschule BFH ◇ TI und • *et* AHB

V.1.12.29 / df / 20. Dezember 2011

Doppelsprachiger Text • *Texte bilingue*

Enthält Teil 6 (Kombinatorik) eines Repetitoriums und Textbuchs zur Begleitung und Ergänzung des Unterrichts. • *Contient partie 6 (analyse combinatoire) d'un cours de répétition et livret à l'accompagnement et complément de la leçon.*
 Produziert mit LaTeX auf NeXT-Computer/ PCTeX WIN98. • *Produit avec LaTeX sur NeXT-Computer/ PCTeX WIN98.*
 Einige Graphiken sind auch mit *Mathematica* entstanden.
 1999 hat der Autor einen Computer-Crash erlebt. Infolge des dadurch provozierten Systemwechsels haben einige Graphiken sehr gelitten. Sie werden neu erstellt, sobald die Zeit dafür vorhanden ist.

• *Quelques graphismes ont été produits avec Mathematica. 1999, l'auteur a été violé d'un crash d'ordinateur. A la suite du changement de système provoqué par cela, quelques graphismes ont souffert tellement. Ils seront aménagés de nouveau dès qu'il y aura assez de temps à disposition pour cela.*

Glück hilft manchmal, Arbeit immer ...

Brahmanenweisheit

• *Des fois le bonheur aidera, mais le travail aide toujours ...*

Sagesse de l'Inde

Aktuelle Adresse des Autors (2007):

Rolf W. Wirz-Depierre

Prof. für Math.

Berner Fachhochschule (BFH), Dep. AHB und TI

Pestalozzistrasse 20

Büro B112 CH-3400 Burgdorf/BE

Tel. ++41 (0)34 426 42 30 / intern 230

Mail: Siehe <http://rowicus.ch/Wir/indexTotalF.html> unter „Koordinaten von R.W.“

Alt: Ingenieurschule Biel (HTL), Ing'schule des Kt. Bern, Fachhochschule ab 1997) // BFH HTA Biel // BFH TI //

©1996, 2001/02/03/04/05/06/07/08/09/10/11

Vor allem die handgefertigten Abbildungen sind früheren öffentlichen Darstellungen des Autors entnommen. Die Urheberrechte dafür gehören dem Autor privat.

• *En particulier les illustrations fait à la main sont pris des représentations publiques et plus précoces de l'auteur. Les copyrights pour cela appartient en privé à l'auteur.*

Die hier noch verwendete LaTeX-Version erlaubt keine Änderung der Überschrift des Inhaltsverzeichnisses . . .

- *La version LaTeX actuellement utilisée n'accepte pas un changement du titre du tableau des matières en français – c'est un programme . . .*

Inhaltsverzeichnis

1 Organisatorisches — Quant à l’organisation	1
2 Grundfragen der Statistik — Questions fondam. de la statist.	3
2.1 Einführung — Introduction	3
2.1.1 Stochastik und Statistik — Stochastique et statistique	3
2.1.2 Das Wesen der Statistik — La nature de la statistique	4
2.1.3 Problemkreise — Domaines de problèmes	6
2.1.4 Arten von Massenerscheinungen — Types d’événements de masse	6
2.1.5 Stichproben — Echantillons	8
2.2 Arbeitsweise der math. Statistik — Façon der travailler d. l. stat. math.	9
2.3 Beschreibende Statistik — Statistique descriptive	10
2.3.1 Fragen — Questions	10
2.3.2 Häufigkeitsverteilungen — Répartition de fréquences	10
2.3.3 Häufigkeitsfunktion — Fonction de fréquences	12
2.3.4 Darstellungstechniken — Techniques de représentation	13
2.3.5 Zur Klassenbildung — Quant à la formation ce classes	17
2.4 Masszahlen einer Stichprobe — Mesures d’un échantillon	18
2.4.1 Mittelwert und empirische Varianz — Moyenne et variance empirique	18
2.4.2 Vereinfachungsmeth. bei Berechnungen — Méthodes de simplif. aux calculs	21
2.4.3 Berechnungen, Häufigkeitsfunktion — Calculs, fonction de fréquence	22
2.4.4 Häufigkeitsverteilung, Massenverteilung — Distr. de fréq. et de masse	23
2.4.5 Beispiel mit Mathematica — Exemple avec Mathematica	24
2.5 Auswertung: Beispiel — Exploitation: Exemple	25
2.5.1 Dateneingabe — Entrée de données	25
2.5.2 Kenngrößen — Caractéristiques	26
2.5.3 Darstellung mittels Kenngrößen — Représentat. par caract.: BoxWhiskerPlot	26
2.5.4 Andere statistische Plots — D’autres représentation statistiques	28
2.6 Weitere Kenngrößen — D’autres caractéristiques	29
2.6.1 Diverse Mittelwerte einer Verteilung — Certaines valeurs moyen. d’une distr.	29
2.6.2 Momente einer Verteilung — Moments d’une distribution	31
2.6.3 Die Schiefe einer Verteilung — Le biais d’une distribution	33
2.6.4 Kurtosis und Exzess — Kurtosis et excès	33
2.6.5 Sinn und Gefahr von Kenngrößen — Caractéristiques: Sens propre et danger	34
3 Kombinatorik — Analyse combinatoire	37
3.1 Einleitung — Introduction	37
3.1.1 Problemstellung — Problème	37
3.1.2 Fakultäten — Factorielles	37
3.2 Anordnungsprobleme — Problèmes d’arrangement	38
3.2.1 Permutationen ohne Wiederholung — Permutations sans répétition	38
3.2.2 Permutationen mit Wiederholung — Permutations avec répétition	42

3.3	Auswahlprobleme — Problèmes de choix	45
3.3.1	Die Fragestellungen — Les questions	45
3.3.2	Variation ohne Wiederholung — Arrangement sans répétition	49
3.3.3	Kombination ohne Wiederholung — Combinaison sans répétition	49
3.3.4	Variation mit Wiederholung — Arrangement avec répétition	52
3.3.5	Kombination mit Wiederholung — Combinaison avec répétition	54
3.4	Übungen — Exercices	56
4	Wahrscheinlichkeitsrechnung — Calcul des probabilités	57
4.1	Einleitung — Introduction	57
4.1.1	Problemstellung — Problème	57
4.1.2	Anwendung — Application	57
4.1.3	Personen — Personnages	58
4.2	Zufallsexperiment, Ereignis — Expérience de hasard, événement	59
4.2.1	Zufallsprozesse, Häufigkeit — Processus aléatoire, fréquence	61
4.3	Ereignisalgebra — Algèbre des événements	64
4.3.1	Ereignisalgebra und Mengenalg. — Alg. des événem. et alg. d. ensembles	64
4.3.2	Boolsche Algebren — Algèbres de Boole	66
4.3.3	Zur Mächtigkeit und Häufigkeit — Quant à la puissance et la fréquence	67
4.3.4	Ereignisbäume — Des arbres d'événements	68
4.4	Klass. Wahrscheinlichkeit n. Laplace — Probabilité class. d'a. Lapl.	68
4.4.1	Laplace-Experiment, Gleichwarsch. — Exp. de Laplace, probab. identique	68
4.4.2	Wahrsch. als Mass für Gewinnchancen — Prob. comme mesure p. gagner	69
4.5	Axiomatischer Wahrsch'begriff — Probabilité axiomatique	71
4.5.1	Begriff, Axiome, Folgerungen — Notion, axiomes, conclusions	71
4.5.2	Der Begriff Wahrscheinlichkeitsraum — La notion espace de probabilité	74
4.5.3	Bedingte Wahrscheinlichkeit — Probabilité conditionnelle	75
4.5.4	Totale Wahrscheinlichkeit — Probabilité totale	79
4.6	Wahrscheinlichkeitsverteilungen — Fonctions de répartition	81
4.6.1	Zufallsvariablen — Variables aléatoires	81
4.6.2	Verteilungsfunktion — Fonction de répartition	84
4.6.3	Verteilungstypen — Types de répartitions	85
4.6.4	Diskrete Verteilung — Répartition discrète	86
4.6.5	Kontinuierliche Verteilung — Répartition continue	90
4.7	Mass- oder Kennzahlen einer Verteilung — Mesures de répartition	92
4.7.1	Allgemeines — Considérations générales	92
4.7.2	Mittelwert — Valeur moyennne	93
4.7.3	Erwartungswert — Valeur d'espérance	94
4.7.4	Symmetrische Verteilung — Distribution symétrique	95
4.7.5	Varianz, Standardabweichung — Variance, écart-type	96
4.7.6	Momente einer Verteilung — Moments d'une distribution	97
4.7.7	Schiefe einer Verteilung — Dissymétrie d'une distribution	100
4.7.8	Weitere Kenngrößen — Autres caractéristiques	100
4.7.9	Momentenerzeugende Funktion — Fonction char. génér. de moments	102
4.7.10	Laplace- und z-Transformation — Transf. de Laplace et en z	103
4.8	Spezielle diskrete Verteilungen — Distributions discrètes spéciales	103
4.8.1	Bernoulliverteilung — Distribution de Bernoulli	103
4.8.2	Gesetze für die Binomialverteilung — Lois pour la distribution de Bernoulli	105
4.8.3	Poissonverteilung — Distribution de Poisson	107
4.8.4	Pascalverteilung — Distribution de Pascal	109
4.8.5	Geometrische Verteilung — Distribution géométrique	110
4.8.6	Hypergeometrische Verteilung — Distribution hypergéométrique	111
4.8.7	Beispiel — Exemple	113

4.9	Spezielle stetige Verteilungen — Distributions continues spéciales	114
4.9.1	Allgemeines — Généralités	114
4.9.2	Rechtecksverteilung — Distribution rectangulaire	114
4.9.3	Normalverteilung — Distribution normale ou de Gauss	115
4.9.4	Grenzwertsätze von Moivre Laplace — Théorèmes limites de Moivre Laplace	118
4.9.5	Lokaler Grenzwertsatz — Théorème limite locale	118
4.9.6	Grenzwertsatz von De Moivre/ Laplace — Théorème limite de De Moivre/ Laplace	120
4.9.7	Das Gesetz der grossen Zahlen von Bernoulli — La loi des gands nombres	120
4.9.8	Bemerkung zum Zufall — Remarques quant au hasard	121
4.9.9	Tschebyscheffsche Ungleichung — Inéquation de Tschebyscheff	122
4.9.10	Logarithmische Normalverteilung — Distribution normale logarithmique	123
4.9.11	Exponentialverteilung — Distribution exponentielle	124
4.9.12	Weibullverteilung — Distribution Weibull	125
4.9.13	Gammaverteilung — Distribution gamma	126
4.9.14	Ausblick — Autres distributions	126
4.10	Zufallsvektoren — Vecteurs aléatoires	126
4.10.1	Fragestellung, Begriffe — Question, notions	126
4.10.2	Der diskrete Fall — Le cas discret	130
4.10.3	Der stetige Fall — Le cas continu	132
4.11	Mehrdimensionale Erwartung — Espérance multidimensionnelle	133
4.11.1	Erwartung, Mittelwert — Espérance, moyenne	133
4.11.2	Varianz, Kovarianz, Korrelation — Variance, covariance, corrélation	134
4.11.3	Der diskrete Fall — Le cas discret	137
4.11.4	Der stetige Fall — Le cas continu	137
4.12	Mehrdimensionale Verteilungen — Répartitions multidimensionnelles	138
4.12.1	Zweidimensionale Normalverteilung — Distribution normale bidimensionnelle	138
4.12.2	Stichprobenfunkt., Testverteilungen — Fonct. d'échant., distrib. de test	139
4.12.3	Chi-Quadrat-Verteilung — Distribution du Khi-deux	143
4.12.4	Sätze zur Chi-Quadrat-Verteilung — Théorèmes sur la distribution du Khi-deux	145
4.12.5	t-Verteilung von Student — Distribution de Student	146
4.12.6	F-Verteilung von Fisher — Distribution de Fisher	148
4.13	Anhang I: Einige Beweise — Annexe I: Certaines preuves	149
4.13.1	Formel zur Gammafunktion — Formule pour la fonction gamma	149
4.13.2	Dichte der Chi-Quadrat-Verteilung — Densité de la distribution Khi-deux	149
4.13.3	Dichte der Student-Verteilung — Densité de la distribution de Student	151
4.13.4	Beweis Tschebyscheffsche Ungleichung — Preuve d'inéquation de Tschebyscheff	151
4.14	Anhang II: Ergänzungen — Annexe II: Suppléments	152
4.14.1	Quadr'summe contra Betrags. — Somme d. carrés com. à la somme d. val. abs.	152
4.14.2	Verteilungstreue u.s.w. — Conformité du type de répartition etc.	153
4.14.3	Zentraler Grenzwertsatz — Théorème limite central	156
4.14.4	Lineare Transformationen — Transformations linéaires	157
5	Math. Statistik — Statist. mathématique	159
5.1	Qualitätskontrolle — Contrôle de qualité	159
5.1.1	Allgemeines, SQC — Généralités, SQC	159
5.2	SQC1: Prozesskontrolle — SQC1: Contrôle de processus	160
5.2.1	Problemstellung — Problème	160
5.2.2	Beispiel Mittelwert — Exemple moyenne	160
5.2.3	Überw'techn. m. Kontr'karten — Techn. de surv. à l'aide d. cartes de contr.	163
5.3	SQC2: Annahmekontrolle — SQC2: Contrôle d'acceptation	163
5.3.1	Hypergeom. Verteil., Urnenmodell — Répart. hypergéom., modèle d'urne	164
5.3.2	Annahmevertrag, Prüfplan — Contrat d'accept., plan d'échant.	165
5.3.3	Binomialmodell, Urnenmodell — Modèle binomial, modèle d'urne	168

5.3.4	Produzenten- u. Konsumentenrisiko — Risque du producteur e.d. consommateur .	169
5.4	Poisson-Prozess, Warteschlange — Files d'attente, processus de Poisson	171
5.4.1	Poisson-Prozess — Processus de Poisson	171
5.4.2	Warteschlangenmodelle — Modèles de files d'attente	175
5.5	Schätzungen — Estimations	176
5.5.1	Schätzer, Punktschätzer — Estimations ponctuelles	176
5.5.2	Erwartungstreue — Propriété d'être sans biais	177
5.5.3	Konsistenz — Consistance	179
5.5.4	Vertrauensintervalle I — Intervalles de confiance I	181
5.5.5	Wichtige Testfunktionen — Fonctions de test importantes	183
5.5.6	Vertrauensintervalle II — Intervalles de confiance II	189
5.5.7	Vertrauensintervalle III — Intervalles de confiance III	193
5.5.8	Vertrauensintervalle IV — Intervalles de confiance IV	196
5.5.9	Automat. Bestimm. der Vertrauensber. — Calcul automat. des domaines de conf. .	196
5.6	Signifikanztests — Tests de signification	197
5.6.1	Hypothesen — Hypothèses	197
5.6.2	Zweiseitige Alternative, t -Test — Alternative bilatérale, test de Student	198
5.6.3	Einseitige Alternative, t -Test — Alternative unilatérale, test de Student	200
5.6.4	Testrisiken — Risques (aléas) aux tests	202
5.6.5	Chi-quadrat-Test für die Varianz — Test khi deux pour la variance	203
5.6.6	Automatisches testen — Tester automatiquement	204
5.6.7	Testrezept im Überblick — Vue d'ensemble d'une recette de test	206
5.6.8	Vergleichstest für Mittelwerte I — Test de compar. de moyennes I	207
5.6.9	Vergleichstest für Mittelwerte II — Test de compar. de moyennes II	209
5.7	Weitere Testverfahren — Autres méthodes de test	211
5.7.1	Eine Typeneinteilung von Tests — Une classification des tests	211
5.7.2	Parameterfreier Test: Vorz'test — Test libre de param.: Test du signe	211
5.7.3	Chi-Quadrat-Anpassungstest — Test d'adaptation de Khi deux	212
5.7.4	Zum Fishertest — Quant au test de Fisher	214
5.7.5	Kontingenztafeln — Tableaux de contingence	215
6	Fehlerech., Regr., Korr. — Calc. de l'err., régr., corr.	219
6.1	Fehlerrechnung (Abhängigkeit) — Calcul de l'erreur (dépendance)	219
6.1.1	Das Problem der Verpflanzung — Le problème de la dépendance	219
6.1.2	Verwendung des totalen Differentials — Appliquer la différentielle totale	220
6.2	Regression — Régression	221
6.2.1	Der Begriff — La notion	221
6.2.2	Methode der kleinsten Quadrate — Méthode des carrés minimaux	222
6.2.3	Korrelation — Corrélation	226
6.2.4	Korrelationskoeff.: Bedeutung — Coeff. de corrélation: Signification	228
6.2.5	Rechnen mit Punktschätzern: Probleme — Calculer avec des estimateurs: Probl. .	229
6.3	Zum Schluss — Quant à la fin	230
A	Aus dem DIYMU: Link	235
B	Fehler von statistischen Kenngrößen und Standardfehler: Link	237
C	Monte-Carlo, Resampling und anderes: Link	239
D	Eine Bootstrap-Anwendung Schritt für Schritt: Link	241
E	Datensatzänderung: Link	243
F	Spezielle Wahrscheinlichkeitssituationen: Link	245

G Hinweise zur Datenanalyse: Link	247
H Crashkurs Kombinatorik, Wahrscheinlichkeit: Link	249

Vorwort zum Teil Kombinatorik • Préface à la partie analyse combinatoire

Probleme mit ganzen Zahlen • *Problèmes aux nombres entiers*

Liebe Leserin, lieber Leser,

Das Thema *Kombinatorik* ist ein klassischer Bestandteil des Mittelschullehrplans. Auch an Berufsmittelschulen sollte es eigentlich behandelt werden. Doch was, wenn ein Student aus irgendwelchen Gründen gerade diesem Stoff an der Schule nie begegnet ist — oder ihn vergessen hat? Dann heisst es eben nacharbeiten und repetieren. Daher ist dieser Text als *Repetitorium* und als *Ausbau* gedacht.

Die Wichtigkeit der Kombinatorik für den Weg durch die weitere Mathematik ist unbestritten. Sie ist ein Werkzeug zur Lösung von Problemen, die manchmal unverhofft an einem herantreten. Geradezu grundlegend ist das Thema aber für das Wissensgebiet „Wahrscheinlichkeitsrechnung und Statistik“.

Dieser Text ist in Skriptform abgefasst. Das bedeutet, dass er in äusserst knapper Fassung nur das wesentliche Skelett des zu lernenden Stoffes wiedergibt. Für weitere, ausführliche Erklärungen, Beispiele, exakte Beweise und ergänzende Ausführungen ergeht daher an den Studenten der Rat, ein oder mehrere Lehrbücher beizuziehen. Studieren bedeutet zu einem wesentlichen Teil, sein Wissen selbständig mit Hilfe der Literatur zu erweitern, streckenweise sogar selbständig zu erarbeiten, zu festigen und anzuwenden. Ein Skript ist dabei nur ein Wegweiser und nie ein Lehrbuchersatz. Welche Lehrbücher jemand verwenden will, ist jedem freigestellt. Das Thema Kombinatorik findet man in praktisch allen Unterrichtswerken für die klassische Gymnasialstufe. Bezüglich der Fachhochschulliteratur sei auf das Beispiel Bibl. A5 (Brenner, Lesky, Mathematik für Ingenieure und Naturwissenschaftler, Band 1) verwiesen.

Im Sommer 1996

Der Autor

• *Chère lectrice, cher lecteur,*

L'analyse combinatoire fait partie du programme du gymnase classique. Dans les écoles qui préparent à la maturité professionnelle, il devrait être traité également. Mais quoi, si un étudiant n'a jamais eu contact avec cette matière pour n'importe quelle raison — ou s'il l'a oubliée? Alors il faut l'élaborer ou répéter. Par conséquent ce texte est conçu comme cours de répétition et comme perfectionnement.

L'importance de l'analyse combinatoire est incontestée. Elle est un outil pour la solution de problèmes qui nous surprennent parfois. Elle est la base pour le "calcul des probabilités et la statistique".

Ce texte est écrit en forme de script. Ça signifie qu'il représente une forme très abrégée de la manière à apprendre. Pour des explications plus vastes et détaillées, exemples, preuves exactes et suppléments, on conseille l'étudiant de consulter plusieurs livres de cours. Etudier signifie en grande partie d'élargir soi-même son savoir à l'aide de la littérature et acquérir de la matière, de l'approfondir et de l'utiliser. Pour cela, un script est seulement un indicateur d'itinéraire et ne remplace jamais un livre de cours. Chacun est libre de choisir ses livres de cours. On trouve le sujet de l'analyse combinatoire pratiquement dans toutes les oeuvres de mathématiques pour le gymnase classique. Concernant le niveau des hautes écoles professionnelles le lecteur est renvoyé à des ouvrages tels que A5 (Brenner, Lesky, Mathematik für Ingenieure und Naturwissenschaftler, Band 1).

Dans l'été 1996

L'auteur

Kapitel 1

Organisatorisches — Quant à l'organisation

Kurze Übersicht • *Vue générale*

1. Organisation, Rahmen • *Organisation, cadre*
2. Stoff • *Matière*
3. Ziel, Weg, Methoden, Feedback, Team • *But, chemin, méthodes, feedback, groupe*
4. Übungen, Selbststudium • *Exercices, études personnelles*
5. Lerntechnik, Arbeitstechnik, Selfmanagement • *Technique de travail et d'apprendre, selfmanagement*
6. Rechte und Pflichten des Studenten und der Schule • *Droits et devoirs de l'étudiant et de l'école*
7. Prinzipien, Grundsätze • *Principes et positions*
8. Rechner, Computer, Mathematiksoftware • *Calculatrices de poche, ordinateurs, software en mathes*
9. Semesterorganisation Mathematik (Anzahl Noten, Prüfungsreglement, Prüfungsplan, Prüfungsrahmen, erlaubte Unterlagen, formale Anforderungen, Benotungskriterien, Benotung der Übungen und Projekte, Arbeitsnachweismappe, Klassensprecher, Klassenbetreuer, Kopierchef, Sprechstunden)
• *Organisation du semestre en mathématiques (nombre de notes, règlement des examens, Plan des examens, conditions concernant les examens, exigences formelles, conditions formelles, philosophie des notes, évaluation des exercices et des projets, porte-feuille, chef de classe, professeur chargé de la classe, chef chargé des copies, heures de consultation)*
10. Hilfsmittel (Bibliothek, Taschenrechner, Mathematiksoftware, Literatur) • *Aides (bibliothèque, calculatrice de poche, software en mathématiques, littérature)*

- 11. Zeitplanung • *Plan du temps à disposition*
- 12. Einführung • *Introduction*

Kapitel 2

Grundfragen der Statistik — Questions fondamentales de la statistique

2.1 Einführung — Introduction

2.1.1 Stochastik — Stochastique et statistique

Stochastik gilt als Oberbegriff. Dieser Name ist vom altgriechischen Wort „stochastike techne“ abgeleitet. Lateinisch bedeutet dies „ars coniectandi“, das heisst „Ratekunst“, Kunst des Vermutens. Die Stochastik ist eigentlich ein Teilgebiet der Mathematik, obwohl man sie heute in eigens für dieses Gebiet eingerichteten Studiengängen lehrt. Man kann also an einigen Universitäten als „Statistiker“ abschliessen, ist dann in Wirklichkeit aber „Stochastiker“. Unter dem Oberbegriff „Stochastik“ werden die Gebiete „Wahrscheinlichkeitstheorie“ und „Statistik“ zusammengefasst.

- **Stochastique** est considérée comme le terme générique. Ce nom est dérivé des mots grecs „stochastike techne“. En latin cela signifie „ars coniectandi“ ce qui signifie „art des taux“, art de la conjecture. La stochastique est proprement une branche des mathématiques, bien qu’aujourd’hui on l’enseigne dans des branches d’étude arrangées exprès pour elle. A quelques universités on peut terminer ainsi les études comme „statisticien“ ce qui est en réalité „stochasticien“. Sous le terme général de „stochastique“ on réunit les domaines scientifiques de la „théorie probabiliste“ et de la „statistiques“.

Unter **Statistik** dagegen versteht man heute die Zusammenfassung bestimmter Methoden zur Analyse von empirischen Daten. Manchmal werden aber die Begriffe Stochastik und Statistik nicht so genau getrennt, da der Begriff „Stochastik“ nicht jedermann bekannt und daher wenig werbewirksam ist. Viele Autoren betiteln ihre Werke mit „Statistik“, meinen damit aber „Stochastik“. Dieser Begriffsgebrauch wird dann in den Werken oft so fortgesetzt.

- *Par contre aujourd’hui on comprend sous „statistique“ le résumé de méthodes pour l’analyse de données empiriques. Parfois les notions „stochastiques“ et „statistiques“ ne sont pas distinguées si exactement, car l’idée de la notion „stochastiques“ n’est pas connue à chacun et par conséquent n’a pas de grande efficacité publicitaire. Beaucoup d’auteurs intitulent leurs livres avec „statistiques“, mais ils pensent à „stochastiques“. Ce mauvais usage persiste souvent dans le texte du livre.*

Das Wort Statistik leitet sich vermutlich vom lateinischen „statisticum“ ab und bedeutet „den Staat betreffend“. Andererseits gibt es in der italienischen Sprache das Wort „statista“, womit Staatsmann resp. Politiker gemeint ist. Im 18. Jahrhundert kam im deutschen Sprachbereich das Wort „Statistik“ in Gebrauch. Man bezeichnete damit die Lehre von den Daten über den Staat (Staatstheorie). Die Wortbedeutung von Statistik als sammeln und auswerten von Daten stammt vermutlich erst aus dem 19. Jahrhun-

dert.

- *Le mot "statistiques" se déduit probablement du mot latin "statisticum" qui signifie "ce qui concerne l'état". D'autre part dans la langue italienne il existe le mot "statista" qui signifie homme d'Etat resp. politicien. Au 18ème siècle dans la langue allemande on a commencé à utiliser le mot "statistique". Par ce mot on spécifiait des données sur l'état (théorie d'état). La signification du mot de statistique comme collectionner et interpréter des données date probablement du 19ème siècle.*

2.1.2 Das Wesen der Statistik — La nature de la statistique

Die grosse Bedeutung der Statistik liegt in ihrer Natur: Die Statistik dient zur Beurteilung von **Massenerscheinungen**. Massenerscheinungen entsprechen unserem Zeitgeist. Neue Methoden machen die Statistik sehr leistungsfähig. Alte empirische Verfahren der Handwerker genügen nicht mehr.

- *La grande importance de la statistique est inhérente à elle-même: La statistique sert à élaborer un jugement sur les événements de masse. Ces événements de masse correspondent à l'esprit de notre siècle. De nouvelles méthodes rendent les statistiques très efficaces. Les vieilles méthodes empiriques des artisans ne suffisent plus.*

Man unterscheidet aktuell zwischen folgenden Arten von Statistik:

1. **Beschreibende, deskriptive oder empirische Statistik**
2. **Schliessende, induktive oder mathematische Statistik, auch Interferenzstatistik**
3. **Hypothesen generierende, explorative oder datenschürfende Statistik, auch Datenschürfung oder Data mining**

- *Actuellement on fait la différence entre les types suivants des statistiques:*

1. *La statistique* **descriptive ou empirique**
2. *La statistique* **inductive ou mathématique**
3. *La statistique* **générant des hypothèses, explorative, extraction des données ou data mining**

Die deskriptive oder beschreibende Statistik ist mathematisch harmlos. Es geht hier um die Darstellung oder Präsentation von statistischem Material.

Bei der mathematischen Statistik (beurteilende, induktive Statistik) will man zu affirmativen Resultaten gelangen. Es geht da um Schätzverfahren und Testverfahren für Hypothesen. Die Resultate sind Wahrscheinlichkeitsaussagen. Hier fliesst also die Wahrscheinlichkeitsrechnung ein. Gesucht werden Zusammenhänge oder Unterschiede in Datenbeständen. Auch sucht man das Mass an Sicherheit von Vermutungen, hypothetischen Ergebnissen oder Modellen. Als **statistisch gesichert** erachtet man ein Resultat erst dann, wenn im Rahmen von vorausschauenden, prospektiven Versuchsplanungen Hypothesen mit einer sehr grossen Wahrscheinlichkeit bestätigt werden können. Zu diskutieren geben hier allemal die Form der Fragestellung in Bezug auf Alternativen sowie das akzeptierte Wahrscheinlichkeitsmass.

Die explorative Statistik (hypothesen-generierende Statistik oder Datenschürfung) ist von den Methoden her eine Zwischenform der beiden vorher genannten Statistikarten. Wegen dem zunehmenden Anwendungsbedarf und ihrer damit wachsenden Bedeutung erleben wir hier zur Zeit (Jahrhundertwende) die Entwicklung einer eigenständigen Anwendungsform.

- *La statistique descriptive est mathématiquement simple ou inoffensive. Il s'agit ici de la présentation ou représentation du matériel statistique.*

Par la statistique mathématique (statistique analytique ou inductive), on veut parvenir aux résultats

affirmatifs. Il s'agit de méthodes d'estimation, d'évaluation ou de méthodes de test pour des hypothèses. Comme résultats on obtient des expressions probabilistes. Ici le calcul de probabilité entre comme partie de la théorie. On cherche des rapports, des relations ou des différences dans les données. On cherche aussi la mesure de la sécurité des suppositions, des résultats hypothétiques des modèles. On estime un résultat comme statistiquement assuré seulement si un résultat peut être confirmé dans le cadre de planifications d'expériences prévoyantes et prospectives avec une probabilité très grande. Ici il faut toujours discuter la forme de la question posée par rapport aux alternatives ainsi que la probabilité acceptée.

Du point de vue des méthodes, la statistique explorative ou générant des hypothèses est une forme intermédiaire des deux méthodes ou des différentes statistiques mentionnées en haut. Comme le besoin d'application et l'importance de la statistique augmentent, nous sommes actuellement témoins (fin de siècle) d'un développement d'une forme d'application indépendante de cette science.

Empirische oder beschreibende Statistik kennen wir schon vom alten Ägypten her (Zwiebelstatistik beim Pyramidenbau ...). Heute geht es aber sehr oft darum, aus einer durch Zahlen beschriebenen Situation durch eine „mathematische Auswertung“ zu einer Beurteilung zu kommen, die mit andern Beurteilungen vergleichbar ist. Die Methoden zur Gewinnung und Auswertung von statistischem Material liefert die **mathematische Statistik**.

• On distingue entre **statistique analytique** ou **statistique descriptive** qui est mathématiquement inoffensive et qui a comme but la représentation ou présentation de matériel statistique (données). Nous distinguons entre la **statistique mathématique** (**statistique analytique** ou **inductive**) qui mène aux résultats affirmatifs ou qui travaille de façon explorative.

Nous connaissons la statistique descriptive et empirique déjà du temps de l'antiquité égyptienne (statistique d'oignons à l'occasion du bâtiment des pyramides ...). Aujourd'hui il s'agit très souvent d'arriver, partant d'une situation décrite par des nombres, à un jugement qui est comparable aux autres jugements, à l'aide d'une "exploitation" mathématique. Les méthodes pour obtenir et exploiter le matériel statistique sont fournies par les **statistiques mathématiques**.

Von Statistiken fordert man folgende fünf Eigenschaften, was leicht nachvollziehbar ist:

1. **Objektivität** (Unabhängigkeit vom Standpunkt des Erstellers der Statistik)
 2. **Reliabilität** (Verlässlichkeit)
 3. **Validität** (Über den Kontext hinausgehende Gültigkeit)
 4. **Signifikanz** (Das Mass an Bedeutung oder des Zutreffens, die Wesentlichkeit)
 5. **Relevanz** (Wichtigkeit)
- Des statistiques, on exige les cinq qualités suivants, ce qui est facilement compréhensible:
1. **Objectivité** (Indépendance du point de vue du producteur de la statistique)
 2. **Reliabilité** (Fiabilité)
 3. **Validité** (Une validité qui dépasse le contexte)
 4. **Portée** (La mesure d'être juste, la mesure d'être essentiel ou fondamental)
 5. **Importance** (La mesure d'être signifiant)

2.1.3 Problemkreise der mathematischen Statistik — Domaines de problèmes de la statistique mathématique

Bezüglich des Vorgehens bei der mathematischen Statistik können wir folgende Problemkreise unterscheiden:

• *Concernant la méthode de la statistique mathématique, nous pouvons distinguer les domaines de problèmes suivants:*

1. Art und Methoden der Datensammlung
 - *Manière et méthodes de la collection des données*
2. Gliederung, Zusammenfassung
 - *classement, structure, disposition, résumé*
3. Analyse: Auswertung, Schlüsse, Hypothesenprüfung
 - *Analyse: Exploitation, conclusion, examen d'hypothèse*
4. Bewertung der Zuverlässigkeit
 - *Estimation la la fiabilité*
5. Beurteilung
 - *Evaluation*
6. Aufzeigung statistischer Gesetzmässigkeiten
 - *Présentation de régularités statistiques*

Anwendung: Überall, wo Massenerscheinungen eine Rolle spielen. Z.B. Reparaturstatistik, Verkaufstatistik, Lager- oder Materialstatistik, Unfallstatistik, Geburten-, Sterbe-, Bevölkerungsstatistik, Gesundheitsstatistik, Klimastatistik, Statistiken in der industriellen Fertigung zwecks Qualitätskontrolle u.s.w..

Der Gewinn aus der Statistik ist die Erarbeitung strategischer Zahlen als Grundlage strategischer Entscheide. Die Absicht ist oft die Risikoverminderung, die Kostensenkung, die Gewinnerhöhung oder die Vermeidung von Verlusten u.s.w..

• **Application:** *Partout où des événements de masse jouent un rôle. P. ex. statistiques de réparation, statistiques de vente, statistiques de stock ou de matériaux, statistiques d'accidents, de naissances-, de décès-, statistiques de population, statistiques de santé, statistiques de climat, statistiques dans la fabrication industrielle en vue de contrôles de qualité u.s.w..*

Le profit qu'on obtient des statistiques est la mise au point de nombres stratégiques comme base de décisions stratégiques. L'intention est souvent la diminution de risques, l'abaissement de frais, l'élévation du gain ou l'évitement de pertes etc..

2.1.4 Arten von Massenerscheinungen — Types d'événements de masse

Man kann zwischen **zwei grundsätzlich verschiedenen Arten von Massenerscheinungen** unterscheiden:

- *On peut distinguer entre deux sortes fondamentalement différentes d'événements de masse:*
1. Grosse Zahl von **Individuen** (Zählstatistiken). Das Zahlenmaterial besteht aus Anzahlen ($\in \mathbb{N}$).
 - *Grand nombre d'individus (statistiques où on compte). Le matériel des données consiste en nombres ($\in \mathbb{N}$).*
 2. Grosse Zahlen von **Einzellerscheinungen** oder meist physikalischen Messungen. Das Zahlenmaterial besteht aus Messwerten, d.h. meist reelle Zahlen ($\in \mathbb{R}$).
 - *Grands nombres d'événements individuels ou généralement de mesurages physiques. Le matériel des données consiste en valeurs indiquées, c.-à.-d. généralement des nombres réels ($\in \mathbb{R}$).*

Beim beobachteten Material (Individuen, Einzelercheinungen) kann es sich allgemeiner um örtliche, zeitliche oder sonst irgendwie geartete unabhängige Dinge handeln, welche wir hier **Beobachtungseinheiten** oder **Merkmale** nennen wollen, die sich auf verschiedene Arten manifestieren oder verschiedene **Ausprägungen** zeigen können. Wir wollen folgende Ausprägungsarten klassifizieren:

- *Quant au matériau observé (individus, phénomènes uniques), il peut s'agir généralement de choses indépendantes et localement et temporellement ou n'importe comment de nature différente, que nous appellons des **unités d'observation** ou des **caractéristiques** qui se manifestent par **sortes ou empreintes différentes**. Nous voulons classifier les sortes de manifestations suivantes:*

1. **Qualitative** Merkmale, d.h. nicht quantifizierbare Merkmale.

Beispiel: Etwas wird als „schön“, „weniger schön“, „unschön“ oder „scheusslich“ qualifiziert. Eine Übertragung auf eine Skala ist hier immer problematisch, da es im vorliegenden Fall um menschliche Gefühle geht, welche nicht gut gegeneinander abgrenzbar und bei verschiedenen Einschätzern nicht unbedingt vergleichbar sind. Dieses Problem taucht oft bei einer Notengebung auf. Mit den hier vorliegenden Ausprägungen zu „rechnen“ scheint unsinnig. Man kann sich ja fragen, ob schön plus scheusslich dann noch schön, weniger schön, unschön oder gar schon scheusslich sein soll. Dasselbe gilt für eine etwaige Multiplikation statt der Addition u.s.w.

- *Les caractéristiques **qualitatives**, c.-à.-d. des caractéristiques non quantifiables.*

Exemple: On qualifie quelque chose de "beau", "moins beau", "disgracieux" ou "horrible". Un transfert sur une échelle graduée est toujours problématique parce qu'il s'agit dans le cas présent de sentiments humains, qu'on ne peut pas bien délimiter l'un contre l'autre et qui ne sont pas absolument comparables encore moins s'il s'agit de taxateurs différents. Ce problème émerge souvent à l'occasion de notes d'école. Il paraît absurde de "calculer" avec des manifestations pareilles. On peut se demander, si "beau" plus "horrible" devrait donc être "beau", "moins beau", "disgracieux" ou "horrible". Le même vaut pour une multiplication éventuelle au lieu de une addition.

2. **Quantitative** Merkmale. Hier kann man zwei Gruppen unterscheiden:

- *Des caractéristiques **quantitatives**. Ici, on peut distinguer deux groupes:*

(a) **Diskret** quantifizierbare Merkmale, z.B. Anzahlen von Individuen. Man denke etwa an die Anzahl Menschen in einem Lift. Hier werden Nachkommastellen absurd, wenn es sich nicht um Kenngrößen wie etwa einen Mittelwert handelt. Wir werden kaum eine Liftkabine finden, in der sich z.B. „4.852 Menschen“ befinden. . .

- *Des caractéristiques quantifiables de façon **discrète**, par exemple des nombres d'individus. On pense peut-être au nombre de gens dans un ascenseur. Ici, les places après le virgule deviennent absurdes, s'il ne s'agit pas de caractéristiques comme par exemple une moyenne. Nous ne trouverons pas une cabine d'ascenseur dans laquelle on observe par exemple "4.852 personnes".*

(b) **Stetig** quantifizierbare Merkmale, z.B. Messresultate, die an einer Skala abgelesen worden sind. So habe man etwa die Länge eines Stabes gemessen. Dieser misst 3.952 m, wobei die letzte Stelle ungenau ist (± 0.001 m). Dabei ist hier zu beachten, dass bei Naturbeobachtungen keine unendlich grosse Genauigkeit möglich ist. „Stetig“ ist also hier in einem weniger exakten, d.h. nicht mathematisch exakten Sinn zu verstehen.

- *Les caractéristiques quantifiables comme **continues**, par exemple des résultats de mesures qu'on a mesurées à l'aide d'une échelle graduée. Imaginons qu'on ait ainsi à mesurer la longueur d'un bâton. Celui-ci mesure 3.952 m à quoi la dernière place est inexacte (± 0.001 m). Quant à cela il faut tenir compte du fait que lorsqu'il s'agit d'observations d'après la nature, il n'est pas possible d'obtenir une exactitude infinie des résultats. Ici, il faut donc comprendre la notion "continue" dans un sens moins exact c.-à.-d., ne le comprendre pas comme l'exactitude mathématique.*

Oft zwingt die Komplexität des Materialumfangs zu einer vorgängigen **Klasseneinteilung** der vorliegenden Zahlen und daher zu einer ersatzweisen Klassendarstellung oder Auswertung.

Im Gegensatz zu den Massenerscheinungen treten bei Individuen oder Einzelercheinungen zufällige Unregelmässigkeiten auf. „Statistische Gesetzmässigkeiten“ fehlen.

• *C'est souvent la complexité de la taille du matériel qui contraint à une **division préalable** des nombres présents en **classes** et par conséquent à une représentation de **classes** de remplacement.*

Contrairement aux événements de masse, il y a des irrégularités accidentelles lors qu'il s'agit d'individus. Les "régularités statistiques" manquent.

2.1.5 Stichproben — Echantillons

Häufig handelt es sich bei den zu untersuchenden Massenerscheinungen um Erscheinungen, die sich aus einer sehr grossen Zahl von Einzelercheinungen zusammensetzen, die man oft unmöglich alle erfassen kann. Daher ist man darauf angewiesen, die Zahl der betrachteten Einzelercheinungen zu reduzieren und auf eine repräsentative Art und Weise eine Auswahl zu treffen. Eine solche repräsentative Auswahl nennen wir **Stichprobe** (Extrakt). Ihr Umfang heisst **Stichprobenumfang**.

• *Souvent il 'agit, lors d'une analyse d'événements de masse, d'événements qui partent d'un très grand nombre d'événements individuels qu'on ne peut absolument pas saisir dans leur totalité. Par conséquent on est forcé de réduire les événements de façon représentative et de faire un choix. Nous appelons un tel choix représentatif un **échantillon** (prise ou récolte d'échantillons). Le nombre de notions individuelles choisies s'appelle **taille d'échantillon**.*

Die erhobenen **Stichprobenwerte** werden in der Regel in einer **Urliste** erfasst (Messprotokoll, Fragebogen u.s.w.). Sie stammen aus einer **Grundgesamtheit** von Ausprägungen, die normalerweise wegen den real vorhandenen Erfassungseinträgen als bekannt oder als im Rahmen der Erfassung bekannt und gegebenenfalls als abgesteckt erweiterbar vorausgesetzt werden kann. So könnte man z.B. bei einer Erfassung einer Biodiversität eine neue Pflanzenart entdecken, die bisher noch in keiner Liste bekannt war. Die Grundgesamtheit wird dadurch jetzt beschränkt erweitert.

• *Les **valeurs d'échantillonnage** recueillies ou rassemblées sont saisies normalement dans une **liste initiale** (protocole des mesurages, questionnaires etc.) Elles sortent d'une **population (univers)** de manifestations, qui normalement, à cause des inscriptions, enregistrements ou mentions réels sont disponibles et connues et, s'il est sensé, sont expansibles de façon délimitée dans le cadre de l'enregistrement. Ainsi on pourrait, à l'occasion d'un enregistrement de la diversité biologique, découvrir par exemple une nouvelle sorte de plante qui n'était jusqu'ici pas connue dans aucune liste. Pour cela, la totalité de la population est élargie maintenant de façon limitée.*

Urlisten kann man natürlich auch gleich ordnen. Dadurch entstehen **geordnete Urlisten**, z.B. eine **Rangliste**. Bei der Erstellung von Ranglisten geht aber schon Information verloren, da dabei die Reihenfolge der Datenerhebung oder der Messungen verändert wird. Diese Reihenfolge kann aber einen systematischen Einfluss auf die Daten haben. Daher liegt das Kapital in der Urliste. Man sollte diese unter allen Umständen als wichtigste Quelle immer aufbewahren, solange an der statistischen Aufbereitung der Daten noch zu arbeiten ist.

• *Evidemment, on peut aussi ordonner une liste initiale tout de suite. Comme ça, on obtient des listes initiales ordonnées, par exemple un classement. Mais lors de la construction d'un classements on perd déjà de l'information, parce que l'ordre de collectionner les données ou de collectionner les mesurages est transformé. Cet ordre peut avoir une influence systématique sur les données. Par conséquent le capital est dans la liste initiale. On devrait la conserver toujours en tous cas comme la source la plus importante tant qu'on travaille encore au traitement statistique des données.*

Eine **Stichprobe** besteht aus **Beobachtungseinheiten** (z.B. Individuen), welche **Merkmale** tragen (Variablen). Diese nehmen in der Regel bei jedem Individuum eine persönliche **Ausprägung** (Wert) an. Die gesammelte Wertemenge nennen wir **Stichprobendaten**.

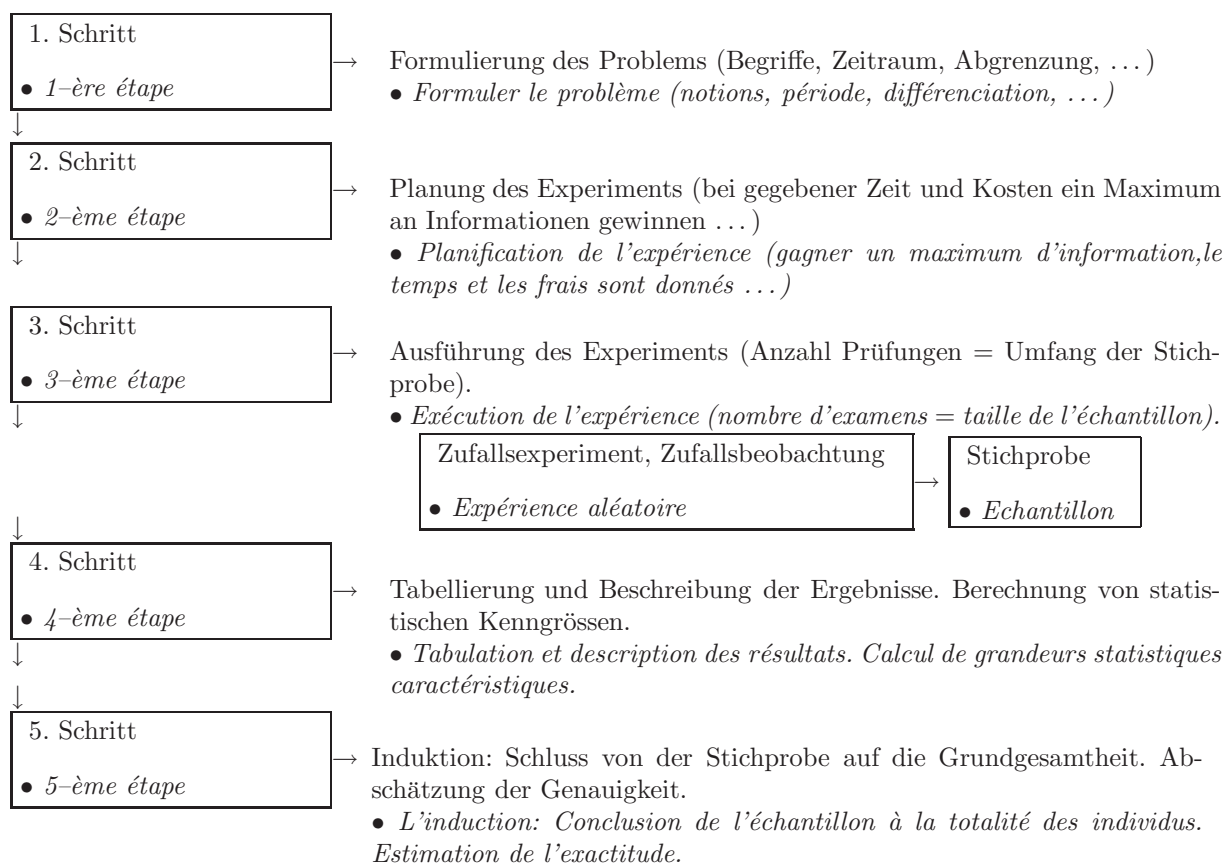
• *Un **échantillon** consiste en **unités d'observation** (par exemple des individus) qui portent des **caractéristiques** (variables). Celles-ci possèdent normalement une **empreinte** ou manifestation person-*

nelle (valeur) pour chaque individu. Nous appelons la quantité des valeurs ramassées les **données d'échantillon**.

2.2 Arbeitsweise der mathematischen Statistik — Façon der travailler dans la statistique mathématique

In der Praxis sind die statistischen Fragestellungen vielfältig. Bearbeitungsschritte und Methoden sind jedoch vielfach gleich:

• *En pratique, les questions statistiques sont variées. Cependant les étapes d'élaboration et les méthodes utilisées sont souvent les mêmes:*



Wichtig: • Important:

1. Die Schlüsse von der Stichprobe auf die Grundgesamtheit fassen auf der Mathematik der Wahrscheinlichkeitsrechnung. Vollkommen sichere Schlüsse dieser Art existieren nicht.
• *Les conclusions de l'échantillon à la totalité des individus se fondent sur les mathématiques du calcul des probabilités. Des déductions de ce genre complètement sûres n'existent pas.*
2. Schlüsse auf die Grundgesamtheit sind nur gültig, wenn wirklich ein **Zufallsexperiment** vorliegt (vgl. 4.2).
• *Les conclusions à la totalité des individus ne sont valables que s'il s'agit vraiment d'une expérience de hasard (voir 4.2).*
3. Ein Experiment kann zur Zerstörung des Artikels führen. Damit wird der Artikel der Zweckbestimmung entzogen.

- Une expérience peut mener à la destruction de l'article. Dans ce cas, l'article est donc soustrait à son affectation.

2.3 Beschreibende Statistik — Statistique descriptive

2.3.1 Fragen — Questions

Stichproben sind in der Regel nur dann sinnvoll, wenn sie aus Zahlen bestehen oder durch Zahlen dargestellt werden können.

- Les échantillons ne sont normalement raisonnables que s'ils consistent en nombres ou s'ils peuvent être représentés par des nombres.

Problem: • **Problème:**

1. Wie können die gesammelten Daten sinnvoll dargestellt werden?
 - Comment est-ce que les données récoltées peuvent être représentées ingénieusement?
2. Wie kann man die Stichproben durch wenige Zahlen kennzeichnen um diese dann so vergleichen zu können? ($\leadsto$ Mittelwert, Varianz ...)
 - Comment est-ce qu'on peut exprimer les échantillons par peu de nombres pour pouvoir les comparer ainsi? ($\leadsto$ Moyenne, variance ...)

Wichtige Begriffe sind **absolute** und **relative Häufigkeit**, **Häufigkeitsfunktion** $\tilde{f}(x)$, **Summenhäufigkeits-** oder **Verteilungsfunktion** $\tilde{F}(x)$ sowie **Kenngrossen** wie etwa ein Mittelwert. Kenngrossen dienen der „Vermessung“ von Daten. Das Ziel ist ihre Darstellung durch wenige Zahlen z.B. zur einfachen Vergleichbarkeit verschiedener Datensätze.

- Les idées importantes sont la **fréquence absolue et relative**, la **fonction de fréquence** $\tilde{f}(x)$, la **fréquence somme (intégrale)** – ou **fonction de répartition** $\tilde{F}(x)$ ainsi que les **caractéristiques** comme par exemple une moyenne. Les caractéristiques servent à mesurer, jauger ou repérer l'ensemble de données. Le but est par exemple la représentation des données par peu de nombres pour obtenir une comparabilité simple de paquets de données différents.

Bemerkung: • **Remarque:** $\tilde{f} \leadsto$ engl. distribution function
 $\tilde{F} \leadsto$ engl. cumulative distribution function on the sample

2.3.2 Häufigkeitsverteilungen — Répartition de fréquences

Die Ergebnisse „Zählung“ oder „Messung“ müssen irgendwie festgehalten oder protokolliert werden. So entsteht eine **Urliste** oder das **Protokoll**. Misst man z.B. an jedem Individuum je zwei Eigenschaften, so hat man schliesslich n Zahlenpaare und daher den Stichprobenumfang n .

- Les résultats d'un échantillonnage (chompter, mesurer) doivent être retenus ou doivent être verbalisés. Ainsi on obtient une liste de base ou le procès-verbal. Si on mesure par exemple à chaque individu deux qualités, on obtient n paires de nombres et par conséquent on a la taille d'échantillon n .

Bsp.: • **Exemple:**

Urliste: Grösse von 200 Pilotenschülern, Messung. Damit die Daten besser darstellbar sind, hat man die Messungen gleich in 40 Klassen (von ca. 150 bis ca. 190) eingeteilt.

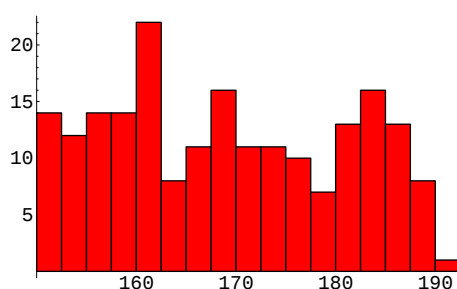
- Liste de base: Grandeur de 200 participants d'un cours de pilotes, mesurer. Afin que les données soient mieux représentables, on a divisé les données directement en 40 classes (d'environ 150 à environ 190).

162.	174.4	182.8	153.2	175.2	162.4	187.6	153.6	184.8	185.6
162.4	178.4	173.6	178.8	167.2	180.	162.	155.6	181.6	186.8
183.2	155.2	150.8	187.6	171.2	170.8	157.6	184.4	186.	158.4
160.4	180.4	151.2	163.2	188.	152.	167.6	174.	170.8	162.
155.6	168.4	179.2	165.2	162.4	163.2	178.4	167.6	180.8	182.
170.8	173.6	184.8	173.6	160.4	182.8	183.6	160.8	162.4	180.8
166.	176.8	181.6	168.8	160.	158.4	152.4	153.6	188.	185.2
164.	175.6	157.2	153.2	183.2	152.4	162.4	169.6	173.2	159.2
168.8	158.8	160.8	168.4	153.2	172.4	168.8	190.	182.8	164.
166.4	186.4	184.8	168.8	152.4	160.4	178.	165.6	158.8	158.4
165.6	186.	176.	188.8	186.8	177.2	165.2	170.4	183.6	154.8
186.4	170.4	150.8	180.8	170.	174.	156.	162.4	167.6	163.6
168.	186.8	158.8	155.2	152.4	150.8	172.8	156.4	155.6	163.6
157.6	176.	162.4	158.8	161.2	155.6	161.6	167.6	181.2	171.6
155.6	155.6	163.6	158.	177.6	158.4	154.8	152.4	175.2	157.6
172.	186.	169.6	184.	164.4	160.	157.2	175.2	153.2	154.4
185.6	157.6	162.	172.8	180.	152.	188.4	165.2	152.4	183.6
183.6	162.4	167.2	175.6	161.6	166.4	187.6	181.6	187.2	156.4
180.4	156.4	183.6	152.	184.4	189.2	171.6	168.8	154.4	187.2
172.8	154.	152.	153.6	179.2	181.6	174.8	168.	167.6	165.2

Darstellung der Stichprobe: • *Présentation de l'échantillonnage*: $\leadsto$ Der Grösse nach geordnete Werte: **Frequenztabelle** (manchmal auch **Strichliste** wie beim Jassen). Die angegebenen Anzahlen heissen **absolute Häufigkeiten**. Z.B. hat der Wert 155.6 die absolute Häufigkeit 6.

• *Valeurs ordonnées d'après la grandeur*: **Tableau de fréquences** (*Des fois aussi liste avec des marques d'après la grandeur, comme les joueurs de cartes*). Les nombres cardinaux donnés s'appellent **fréquences absolues**. P.ex. la valeur 155.6 a la fréquence absolue 6.

(150.8, 3)	(155.6, 6)	(160.4, 3)	(165.2, 4)	(170., 1)	(174., 2)	(178.4, 2)	(183.2, 2)
(151.2, 1)	(156., 1)	(160.8, 2)	(165.6, 2)	(170.4, 2)	(174.4, 1)	(178.8, 1)	(183.6, 5)
(152., 4)	(156.4, 3)	(161.2, 1)	(166., 1)	(170.8, 3)	(174.8, 1)	(179.2, 2)	(184., 1)
(152.4, 6)	(157.2, 2)	(161.6, 2)	(166.4, 2)	(171.2, 1)	(175.2, 3)	(180., 2)	(184.4, 2)
(153.2, 4)	(157.6, 4)	(162., 4)	(167.2, 2)	(171.6, 2)	(175.6, 2)	(180.4, 2)	(184.8, 3)
(153.6, 3)	(158., 1)	(162.4, 8)	(167.6, 5)	(172., 1)	(176., 2)	(180.8, 3)	(185.2, 1)
(154., 1)	(158.4, 4)	(163.2, 2)	(168., 2)	(172.4, 1)	(176.8, 1)	(181.2, 1)	(185.6, 2)
(154.4, 2)	(158.8, 4)	(163.6, 3)	(168.4, 2)	(172.8, 3)	(177.2, 2)	(181.6, 4)	(186., 3)
(154.8, 2)	(159.2, 1)	(164., 2)	(168.8, 5)	(173.2, 1)	(177.6, 1)	(182., 1)	(186.4, 2)
(155.2, 2)	(160., 2)	(164.4, 1)	(169.6, 2)	(173.6, 3)	(178., 1)	(182.8, 3)	(186.8, 3)



Statt eine Frequenz- oder Strichliste zu machen, kann man die Daten auch gleich in ein Histogramm oder in ein Punktdiagramm eintragen. $\leadsto$ Punkt pro Messung, Punkte statt Balken.

• *Au lieu de faire une liste de fréquence ou une liste de marques, on peut inscrire les données tout de suite dans un histogramme ou dans un diagramme de points*. $\leadsto$ Un point par mesure, points au lieu de bâtons.

Wenn man die Häufigkeit als Bruchteil des ganzen Stichprobenumfangs bestimmt, spricht man von der **relativen Häufigkeit**. Durch diese Grösse lassen sich Stichproben mit verschiedenem Umfang vergleichen.

• *Si on calcule la fréquence comme la fraction de la taille de l'échantillon plein, on parle de la fréquence relative*. Par cette grandeur on peut comparer des échantillons d'une taille différente.

Definition: • **Définition:**

Sei $H(a)$ die absolute Häufigkeit der Beobachtung A mit dem Wert $x = a$ und n der Stichprobenumfang.

• Soit $H(a)$ la fréquence absolue de l'observation A dont la valeur soit $x = a$ et n la taille de l'échantillon.

$\leadsto$ **relative Häufigkeit** • **Fréquence relative**

$$h(a) = \frac{H(a)}{n}$$

Es gilt selbstverständlich: • Il vaut clairement:

Satz: • **Théorème:** $0 \leq h(a) \leq 1$

2.3.3 Häufigkeitsfunktion — Fonction de fréquences

Geg.: • **Donné:**

m verschiedene Werte $\{x_1, x_2, x_3, \dots, x_m\}$ mit den relativen Häufigkeiten $h(x_1), h(x_2), \dots, h(x_m)$, kurz $h_1, h_2, \dots, h_m$. Die Werte fallen also in m Klassen $\{[x]_1, \dots, [x]_m\}$. Der Stichprobenumfang sei n .

• m valeurs différentes $\{x_1, x_2, x_3, \dots, x_m\}$ avec les fréquences relatives $h(x_1), h(x_2), \dots, h(x_m)$, bref $h_1, h_2, \dots, h_m$. Les valeurs tombent dans m classes $\{[x]_1, \dots, [x]_m\}$. Le taille d'échantillon soit n .

Definition: • **Définition:****Häufigkeitsfunktion:** • **Fonction de fréquence:**

$$\tilde{f}(x) = \begin{cases} h_i & x = x_i \\ 0 & x \neq x_i \end{cases}$$

Definition: • **Définition:****Verteilungsfunktion:** (Diskrete Verteilung)• **Fonction de distribution:** (Distribution discrète)

$$\tilde{F}(x) = \sum_{x_i \leq x} f(x)$$

Die Häufigkeitsfunktion der Stichprobe zeigt, wie die Häufigkeiten verteilt sind. Sie bestimmen also die **Häufigkeitsverteilung**. Folgende Beziehungen sind unmittelbar einsichtig:

• La fonction de fréquence de l'échantillon montre, comme les fréquences sont distribuées. Elle détermine ainsi la distribution de fréquence. Les relations suivant les sont directement compréhensibles:

Satz: • **Théorème:**

1. $\sum_{i=1}^m H(x_i) = \sum_{i=1}^m H_i \leq \sum_{i=1}^n H(x_i) = n$
2. $\sum_{i=1}^m h(x_i) = \sum_{i=1}^m h_i = \sum_{i=1}^m \tilde{f}(x) = \tilde{F}(x_m) \leq 1 = \tilde{F}(x_n)$

2.3.4 Darstellungstechniken — Techniques de représentation

Beispiel einer Datenmenge — Exemple d'un ensemble de données

Geg.: • **Donné:**

Resultate einer Zählung von defekten Komponenten an 40 Apparaten eines gegebenen Typs, die zur Reparatur eingeschickt worden sind. Notiert sind die Anzahlen x_i der Komponenten.

• *Les résultats d'un comptage de composantes défectueuses de 40 appareils d'un type donné, qui ont été envoyés à la réparation. On a noté les nombres x_i de composantes.*

0, 0, 0, 1, 1, 2, 1, 5, 5, 6,
3, 4, 2, 6, 8, 9, 7, 3, 4, 1,
2, 3, 5, 7, 5, 3, 2, 4, 6, 8,
9, 7, 5, 6, 4, 2, 2, 2, 1, .

Bemerkung: • **Remarque:**

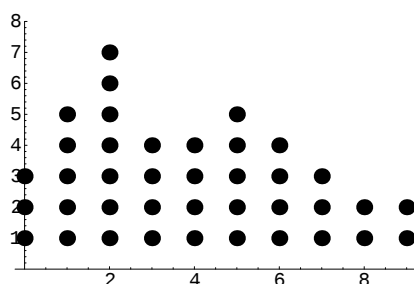
Der Punkt (.) bedeutet **missing value**: Der Wert wurde nicht notiert. Damit wird hier $n = 39$.

• *Le point (.) signifie **missing value**: La valeur n'a pas été enregistrée. Donc $n = 39$.*

↪ Frequenztafel: • *Tableau de fréquences:*

x_i	0	1	2	3	4	5	6	7	8	9	10
H_i	3	5	7	4	4	5	4	3	2	2	—

Punktdiagramm — Diagramme de points



Punktdiagramm (statt Strichliste)

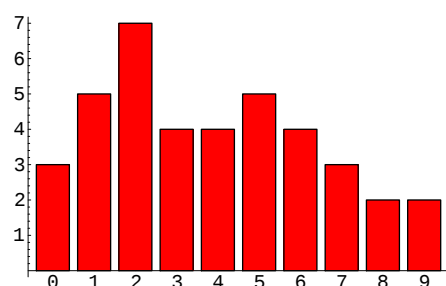
• *Diagramme de points (à la place d'une liste de fréquences)*

Bemerkung: • **Remarque:**

Statt mit Strichlisten kann man auch mit **Stamm-Blatt-Diagrammen** arbeiten. Dabei werden z. B. bei Messungen die vorkommenden Ziffern ohne die letzte in eine Rangliste (Spalte) eingetragen. Die letzten Ziffern werden statt einem Strich wie bei einer Strichliste dahinter in die zugehörige Zeile geschrieben. Das macht Sinn, wenn die Messungen sich häufig nur in der letzten Ziffer unterscheiden.

• *Au lieu de travailler avec des listes de contrôle (avec des traits) on peut aussi travailler avec des **diagrammes de tronc-feuille**. Ici, pour un mesurage, les chiffres, sans le dernier, sont inscrits par exemple dans un classement (colonnes). Les derniers chiffres sont écrits dans la ligne affiliée: Comme à une liste de traits on inscrit le chiffre au lieu d'un trait. Cela a un sens, si les mesurages se distinguent souvent seulement par le dernier chiffre.*

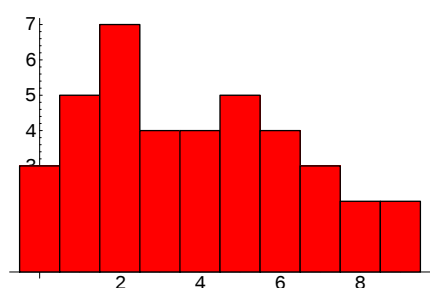
Stabdiagramm — Diagramme en bâtons



Stabdiagramm oder Balkendiagramm

- *Diagramme en bâtons*

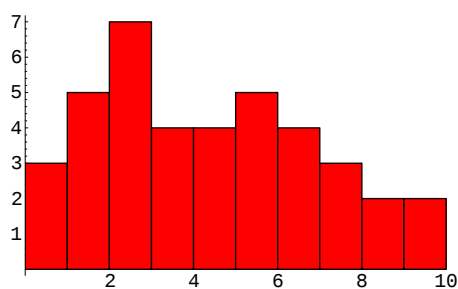
Generalisiertes Stabdiagramm — Diagramme en bâtons généralisé



Generalisiertes Stabdiagramm (hier Balkenbreite grösser)

- *Diagramme en bâtons généralisé (ici la largeur des bâtons est plus grande)*

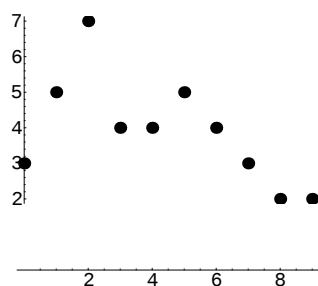
Histogramm — Histogramme



Histogramm oder Staffelfeld (Treppenfunktion), in der Regel bei einer Zusammenfassung der Daten in Klassen

- *Histogramme (fonction en escaliers), normalement pour résumer des données dans des classes*

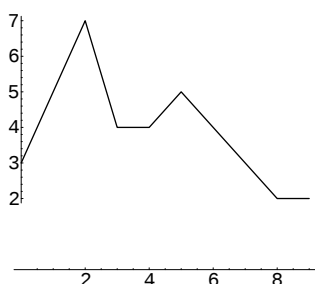
ListPlot — ListPlot



Plot einer Liste

- *Plot d'une liste*

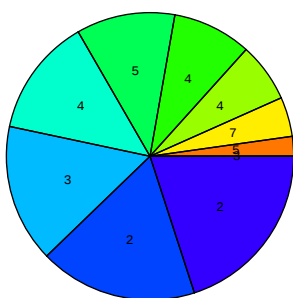
Häufigkeitspolygon — Polygone de fréquences



Häufigkeitspolygon

- *Polygone de fréquences*

Kreisdiagramm — Diagramme à secteurs



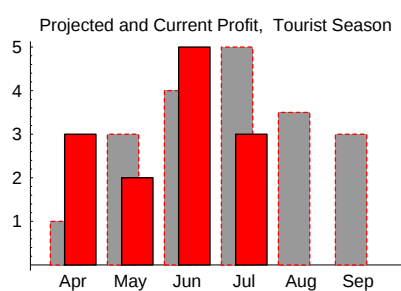
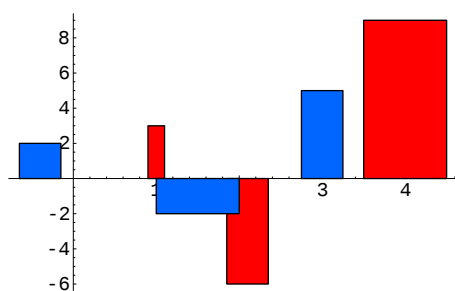
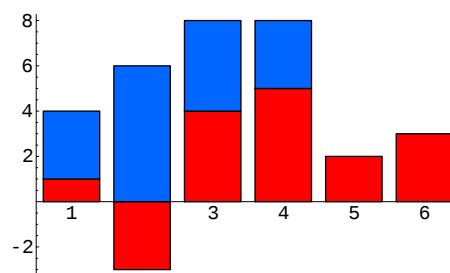
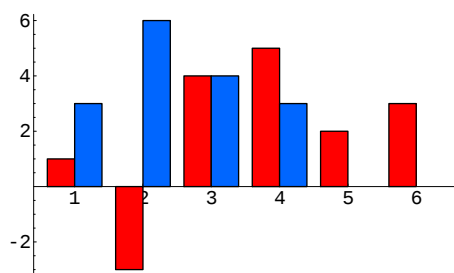
Kreisdiagramm

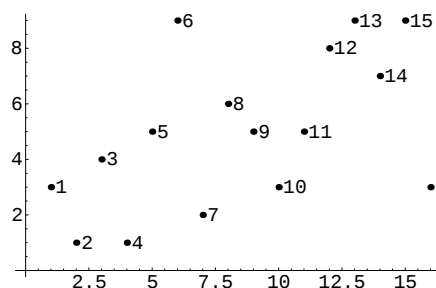
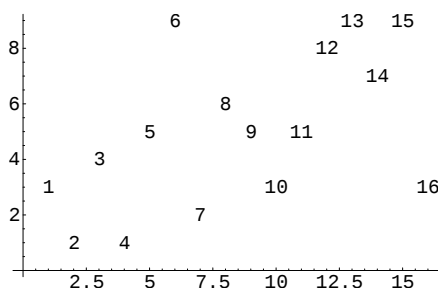
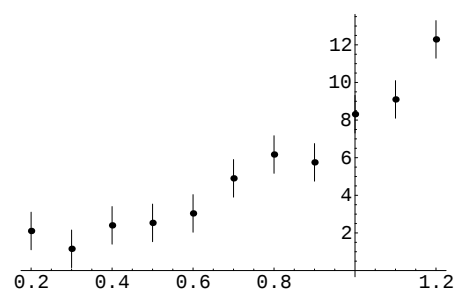
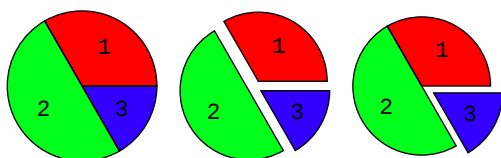
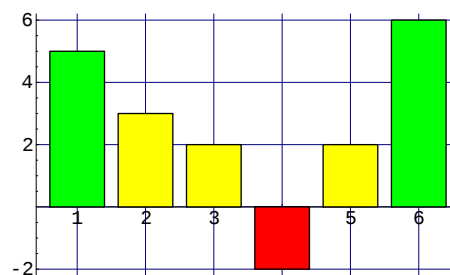
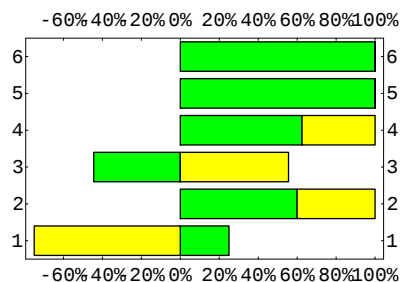
- *Diagramme à secteurs, diagramme en camembert, graphique sectoriel*

Gegliederte Diagramme — Diagrammes structurés

Den folgenden Beispielen liegen Datenmengen zugrunde, die von obigen Daten abweichen. Die Natur der Sache macht dies notwendig. Die Beispiele sind selbsterklärend.

- *Dans l'exemple suivant les quantités de données sont différentes de la quantité de données qu'on vient de traiter. La nature de la chose le rend nécessaire. Les exemples sont évidents.*



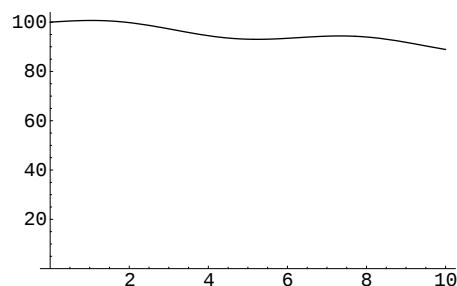
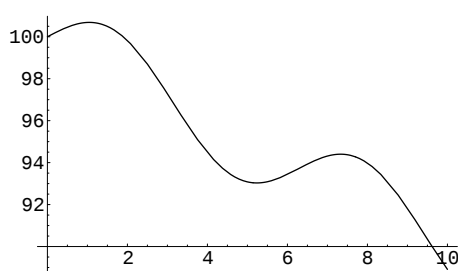


In technischen und ökonomischen Publikationen findet man eine Vielfalt weiterer Darstellungsarten. Erwähnenswert sind **Sankey-Diagramme**, die an Flussdiagramme erinnern und **Körperschaubilder** oder **3D-Darstellungen**.

• Dans les publications techniques et économiques, on trouve une quantité de types de représentations très vastes. P.ex. les **diagrammes de Sankey** qui rappellent des organigrammes ou des schémas de flux. En outre il faut mentionner les **diagrammes tridimensionnels** ou **représentations 3D**.

Achtung Schwindel — Attention escroquerie

Folgendes Beispiel spricht für sich: • L'exemple suivant explique tout:



2.3.5 Zur Klassenbildung — Quant à la formation de classes

Bei grösseren Datenmengen wird eine Tabelle oder Graphik rasch unübersichtlich. Durch Gruppierung der Daten in Klassen (**Klassenbildung**) kann man die Übersicht wieder gewinnen und die Sache vereinfachen. Der gemachte Fehler ist höchstens von der Grösse der halben Klassenbreite.

• *Si on a une quantité de données assez grande, un tableau ou un graphisme deviennent vite peu clairs. Par groupement des données en classes (**formation de classes**), on peut de nouveau s'y retrouver et simplifier la chose. L'erreur faite est au maximum de la grandeur de la moitié de la largeur des classes.*

Dabei wird das Intervall, in dem alle Stichprobenwerte liegen, in **Klassenintervalle** (Teilintervalle) unterteilt. Die Mitten dieser Intervalle heissen **Klassenmitten**. Die in einem Teilintervall liegenden Werte der Stichprobe ist eine Teilstichprobe und heisst **Klasse von Werten**. Die Anzahl der Werte, die in einer Klasse liegen, nennen wir die **zur Klasse gehörige absolute Häufigkeit** resp. **Klassenhäufigkeit**. Dazu gehört:

• *L'intervalle, dans lequel se trouvent toutes les valeurs d'échantillon, soit subdivisé en **intervalles de classes** (intervalles partiels). Les centres de ces intervalles s'appellent **centres de classes**. Le nombre de valeurs qui sont situées dans une classe, s'appellent **fréquences absolues de la classe** ou **fréquences de la classe**. Quant à cela nous retenons:*

Definition: • **Définition:** **Relative Klassenhäufigkeit** $:= \frac{\text{absolute Klassenhäufigkeit}}{\text{Stichprobenumfang}}$
 • **Fréquence de classe relative** $:= \frac{\text{fréquence de classe absolue}}{\text{taille de l'échantillon}}$

Bemerkung: • **Remarque:** Wir nehmen jetzt an, dass alle Werte einer Klasse in der Klassenmitte liegen. (Diese Annahme wird meistens so gemacht.) Das erzeugt einen Fehler, der aber höchstens so gross wie die halbe Klassenbreite ist.

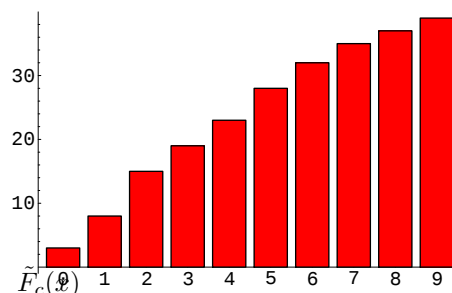
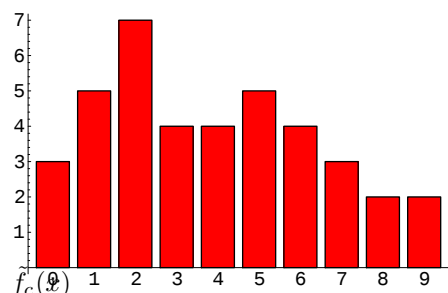
• *Maintenant nous supposons que toutes les valeurs d'une classe soient situées dans le centre de la classe. (Cette supposition est normalement faite de cette manière.) Ça nous produit une erreur qui est au maximum aussi grande que la demi-largeur de la classe.*

Die Häufigkeit in Abhängigkeit von der Klassenmitte wird beschrieben durch die Häufigkeitsfunktion $\tilde{f}_c(x)$ der nun neu in Klassen eingeteilten Stichprobe. Dazu gehört die Summenhäufigkeitsfunktion oder Verteilungsfunktion $\tilde{F}_c(x)$. $\tilde{F}_c(x)$ ist eine Treppenfunktion, die monoton wächst.

• *La fréquence en dépendance du centre de la classe est maintenant décrite par la fonction de fréquence $\tilde{f}_c(x)$ de l'échantillon qui est maintenant divisé en classes. Liée avec $\tilde{f}_c(x)$ est la fonction de fréquence de somme ou fonction de distribution $\tilde{F}_c(x)$. $\tilde{F}_c(x)$ est une fonction en escalier qui croît monotonement.*

Beispiel: Die folgenden Diagramme zeigen $n\tilde{f}_c(x)$ und $n\tilde{F}_c(x)$ dargestellt durch Stabdiagramme. Diese Darstellung ist zwar nicht ganz korrekt, dafür heute aber oft leicht auf dem Computer erzeugbar.

• **L'exemple:** *Les diagrammes suivants montrent $n\tilde{f}_c(x)$ et $n\tilde{F}_c(x)$ représentés par des diagrammes de bâton. En effet cette représentation n'est pas tout à fait correcte, mais elle est aujourd'hui souvent produite facilement par les moyens de l'ordinateur.*



Es ist unmittelbar klar:

- Il est évident qu'il vaut:

Satz: • **Théorème:** $x_1 \leq x_2 \Rightarrow 0 \leq \tilde{F}_c(x_1) \leq \tilde{F}_c(x_2) \leq 1$

Bemerkung: • **Remarque:**

1. $\tilde{f}_c(x)$ und $\tilde{F}_c(x)$ sind konstant ausser für $x = \text{Klassengrenze}$.
 • $\tilde{f}_c(x)$ et $\tilde{F}_c(x)$ sont constantes sauf pour $x = \text{limite de classe}$.
2. $\tilde{f}_c(x)$ und $\tilde{F}_c(x)$ springen an den Klassengrenzen. Die Sprunghöhe von $\tilde{F}_c(x)$ ist gleich $\tilde{f}_c(x)$.
 • $\tilde{f}_c(x)$ et $\tilde{F}_c(x)$ ont des sauts à la limite des classes. La hauteur du saut de $\tilde{F}_c(x)$ est égale à $\tilde{f}_c(x)$.

2.4 Masszahlen einer Stichprobe — Mesures d'un échantillon

Eine Stichprobe kann einerseits durch die Häufigkeitsfunktion oder durch die Verteilungsfunktion beschrieben werden. Andererseits kann man sie auch durch gewisse Masszahlen charakterisieren. Diese Masszahlen sind dann mit den Masszahlen anderer Stichproben zahlenmässig vergleichbar. Das ist ein Vorteil. Der Nachteil ist, dass bei den Masszahlen viel Information nicht mehr direkt sichtbar ist, welche ursprünglich noch da war. Wichtige Masszahlen sind der Mittelwert und die Varianz.

Solche Masszahlen teilt man in verschiedene Kategorien ein. So ist der Mittelwert oder auch das Minimum und das Maximum einer Stichprobe ein **Lagemass**. Die Streuung dagegen ist ein **Streumass** u.s.w.

• *Un échantillon peut être décrit d'une part par la fonction de fréquence ou par la fonction de distribution. D'autre part on peut aussi le caractériser par certains nombres ou certaines mesures. Ces mesures sont numériquement comparables aux mesures d'autres échantillons. C'est un avantage. Le désavantage est que beaucoup d'informations qui étaient au début encore là, ne sont plus directement visibles par les mesures. Les mesures importantes sont la moyenne et la variance.*

*On divise de telles mesures en des catégories différentes. Ainsi la moyenne ou bien le minimum et le maximum d'un échantillon sont des **mesures de position**. Par contre la dispersion ou la variance sont des **mesures de dispersion**.*

2.4.1 Mittelwert und empirische Varianz — Moyenne et variance empirique

Mittelwert — Moyenne

Der Mittelwert zeigt die durchschnittliche Grösse der Stichprobenwerte oder Klassen an. Er ist das arithmetische Mittel der Stichprobenwerte oder Klassenwerte:

- La moyenne montre la grandeur moyenne des valeurs d'échantillon ou des classes. C'est la moyenne arithmétique des valeurs d'échantillon ou des valeurs des classes:

Definition: • **Définition:**

Mittelwert: • **Moyenne**

$$\bar{x} = \frac{x_1 + x_2 + \dots + x_n}{n} = \frac{1}{n} \cdot \sum_{i=1}^n x_i$$

Problem: • **Problème:** Die Stichproben 1, 2, 3, 4, 5 und 2.7, 3.0, 3.1, 3.2 haben denselben Mittelwert, obwohl sie wesentlich verschieden sind. Gesucht ist also eine Masszahl, die die Abweichungen der Stichprobenwerte vom Mittelwert angibt.

- Les échantillons 1, 2, 3, 4, 5 et 2.7, 3.0, 3.1, 3.2 ont la même moyenne, tandis qu'ils sont essentiellement différents. On cherche donc une mesure qui représente la différence des valeurs d'échantillon de la moyenne

Der Mittelwert ist nicht **robust** (wenig empfindlich) gegen **Ausreisser** (z.B. sonderbare, unerklärlich grosse oder kleine Stichprobenwerte). Ein Ausreisser kann den Mittelwert sehr verfälschen.

- La moyenne n'est pas **robuste** (peu sensible), contre des **valeurs aberrantes** (valeurs d'échantillon par exemple étranges, inexplicablement grandes ou petites). Une valeur aberrante peut beaucoup falsifier la moyenne.

Statt dem gewöhnlichen arithmetischen Mittel sind auch andere Mittelwerte in Gebrauch. Solche anderen Mittelwerte gibt es viele: Z.B. das **geometrische**, das **harmonische** oder das **gewichtete arithmetische Mittel** (entsprechend einer gemittelten Momentensumme). Dazu konsultiere man die Literatur.

- Au lieu de la moyenne arithmétique ordinaire on utilise aussi d'autres moyennes. Il y en a beaucoup: Par exemple la **moyenne géométrique**, la **moyenne harmonique** ou le **moyen arithmétique pondéré** (correspondant à une somme de moments pondérés). Pour cela on consulte la littérature.

Empirische Varianz — Variance empirique

Eine Möglichkeit wäre es, die Summe der Differenzen $\sum(x_i - \bar{x})$ zu nehmen. Das funktioniert aber nicht, wie man sofort nachrechnet:

- Une possibilité serait de prendre la somme des différences $\sum(x_i - \bar{x})$. Mais ceci ne fonctionne pas comme on voit tout de suite:

$$\sum_{i=1}^n (x_i - \bar{x}) = \left(\sum_{i=1}^n x_i \right) - (n \cdot \bar{x}) = n \cdot \frac{\sum_{i=1}^n x_i}{n} - n \cdot \bar{x} = n \cdot \bar{x} - n \cdot \bar{x} = 0$$

$x_i - \bar{x}$ kann negativ oder positiv oder auch null werden. In der Summe heben sich die Werte dann auf. Das könnte man umgehen, indem man die Beträge $|x_i - \bar{x}|$ nimmt. Mathematisch einfacher und gewinnbringender ist es aber, statt dessen die Quadrate zu nehmen. In Anlehnung an den Mittelwert der so denkbaren quadratischen Abweichung der Stichprobenwerte x_i vom Stichprobenmittelwert $\bar{x}$ definieren wir:

- $x_i - \bar{x}$ peut être négatif ou positif ou aussi zéro. Mais alors, les valeurs s'annulent dans la somme. Ceci pourrait être évité en prenant les valeurs absolues $|x_i - \bar{x}|$. Mais il est mathématiquement plus simple et plus profitable d'en prendre à cette place les carrés. En nous appuyant à la moyenne de l'écart carré des valeurs d'échantillon x_i de la moyenne d'échantillon $\bar{x}$ ainsi imaginable, nous définissons:

Definition: • **Définition:**

$$s^2 = \frac{1}{n-1} \cdot \sum_{i=1}^n (x_i - \bar{x})^2$$

heisst **Varianz** der Stichprobe $\{x_1, \dots, x_n\}$.• *s'appelle variance de l'échantillon $\{x_1, \dots, x_n\}$.* $s = \sqrt{s^2}$ heisst **Standardabweichung**.• *$s = \sqrt{s^2}$ s'appelle écart-type (étalon, quadratique moyen, déviation normale)*Varianz und Standardabweichung werden auch als **Streuung** bezeichnet.• *(Problème de notions de langue allemande.)***Bemerkung:** • **Remarque:**

Interessanterweise steht bei der empirischen Varianz im Nenner $n - 1$ und nicht n , wie das bei einem Mittelwert zu erwarten wäre. Einerseits hat das bei grossen n praktisch keinen Einfluss auf den Wert von s^2 . Andererseits aber erreicht man damit, dass s^2 eine „erwartungstreue Schätzung der Streuung σ^2 “ der Grundgesamtheit ist. (Vgl. dazu die Ausführungen in Storm, Bibl. A12.)

• *Il est intéressant de remarquer que la variance empirique a dans le dénominateur $n - 1$ et non n , comme on s'attendrait pour la moyenne. D'une part ça n'a pratiquement aucune influence sur la valeur de s^2 si n est assez grand. D'autre part on atteint par cela que s^2 est une "appréciation de la diffusion σ^2 qui correspond à l'espérance mathématique" de l'ensemble de base de données. (Quant à cela voir les exposés dans Storm, Bibl. A12.)*

Die Varianz und die Standardabweichung sind nicht robust gegen Ausreisser.

• *La variance et l'écart standard ne sont pas robustes contre des valeurs aberrantes.***1. Beispiel:** • **Exemple 1:**

$$M_1 = \{1, 2, 3, 4, 5\} \Rightarrow \bar{x} = 3.0, s^2 \approx 2.5, s \approx 1.6$$

2. Beispiel: • **Exemple 2:**

$$M_1 = \{2.7, 3.0, 3.1, 3.2\} \Rightarrow \bar{x} = 3.0, s^2 \approx 0.05, s \approx 0.22$$

Definition: • **Définition:** $[\bar{x} - s, \bar{x} + s]$ heisst **Standardintervall**• $[\bar{x} - s, \bar{x} + s]$ s'appelle **intervalle standard**Im letzten Beispiel ist das Standardintervall ca. $[2.78, 3.22]$.• *Dans l'exemple dernier l'intervalle standard est environ $[2.78, 3.22]$.*

Weitere Masszahlen — Autres mesures

Definition: • **Définition:****Spannweite** der Stichprobe := grösster minus kleinster Stichprobenwert.• **Etendue** de l'échantillon := valeur maximale moins valeur minimale de l'échantillon.**Median** ($\tilde{x}$, Zentralwert) der Stichprobe := Wert in der Mitte bei einer ungeraden Anzahl von Stichprobenwerten, ansonst arithmetisches Mittel der beiden Werte in der Mitte. (Der Median ist robust gegen Ausreisser!)• **Médian** ($\tilde{x}$, valeur centrale) de l'échantillon := valeur au centre si le nombre des valeurs de l'échantillon est impair, si non moyenne arithmétique de la valeur moyenne des deux valeurs au centre. (La valeur centrale (médiane) est robuste contre les valeurs aberrantes!)**Quartile, Quantile:** Vgl. Box-Whisker-Plot, Seite 26.• **Quartile, Quantile:** Voir Box-Whisker-Plot, page 26.**Modus** (D, Dichtemittel) der Stichprobe := Wert mit der grössten Häufigkeit (wird manchmal auch lokal genommen).• **Mode** (D, moyenne de densité) de l'échantillon := valeur à la fréquence maximale. (Des fois, le mode est pris aussi de façon locale.)

Schiefe, Kurtosis und diverse Mittelwerte vgl. Seite 29.

• *Dissymétrie, kurtosis et des valeurs moyennes différentes, voir page 29.*

2.4.2 Vereinfachungsmethoden bei Berechnungen — Méthodes de simplification aux calculs

Es gilt: • *Il vaut:*

$$\begin{aligned}
s^2 &= \frac{1}{n-1} \cdot \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \cdot \sum_{i=1}^n (x_i^2 - 2x_i\bar{x} + \bar{x}^2) = \frac{1}{n-1} \cdot \left(\sum_{i=1}^n x_i^2 - 2\bar{x} \sum_{i=1}^n x_i + \sum_{i=1}^n \bar{x}^2 \right) = \\
&= \frac{1}{n-1} \cdot \left(\sum_{i=1}^n x_i^2 - 2\bar{x}n\bar{x} + n\bar{x}^2 \right) = \frac{1}{n-1} \cdot \left(\left(\sum_{i=1}^n x_i^2 \right) - n\bar{x}^2 \right) = \frac{1}{n-1} \cdot \sum_{i=1}^n x_i^2 - \frac{n}{n-1} \bar{x}^2 = \\
&= \frac{1}{n-1} \cdot \sum_{i=1}^n x_i^2 - \frac{1}{n(n-1)} \left(\sum_{i=1}^n x_i \right)^2 = \frac{1}{n-1} \cdot \left(\sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x_i \right)^2 \right)
\end{aligned}$$

Satz: • **Théorème:**
$$s^2 = \left(\frac{1}{n-1} \cdot \sum_{i=1}^n x_i^2 \right) - \left(\frac{n}{n-1} \cdot \bar{x}^2 \right) = \frac{1}{n-1} \cdot \left(\sum_{i=1}^n x_i^2 - \frac{1}{n} \cdot \left(\sum_{i=1}^n x_i \right)^2 \right)$$

Bei grossen Zahlen x_i ist eine Koordinatentransformation $x_i = c + x_i^*$ vorteilhaft. (c ist so gewählt, dass die x_i^* kleine Zahlen werden.)• *Pour des nombres grands x_i , une transformation de coordonnées $x_i = c + x_i^*$ est favorable. (c choisi de manière que les x_i^* deviennent des nombres petits.)*

$$\leadsto x_i^* = c - x_i$$

$$\bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i = \frac{1}{n} \cdot \sum_{i=1}^n (c + x_i^*) = \frac{1}{n} \cdot \sum_{i=1}^n x_i^* + \frac{1}{n} \cdot \sum_{i=1}^n c = \frac{1}{n} \cdot \sum_{i=1}^n x_i^* + \frac{1}{n} \cdot \sum_{i=1}^n c = \bar{x}^* + \frac{n \cdot c}{n} = \bar{x}^* + c$$

$$s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \sum_{i=1}^n (c + x_i^* - \bar{x}^* - c)^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i^* - \bar{x}^*)^2 = (s^*)^2 \Rightarrow s^2 = (s^*)^2$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$x_i = c + x_i^*$$

Beh.: • **Thè.:**

$$1. \bar{x} = \bar{x}^* + c$$

$$2. s^2 = (s^*)^2$$

$$3. s = s^*$$

Manchmal ist sogar eine **lineare Transformation** nützlich:

• *Des fois une transformation linéaire est utile:*

$$\text{Sei } \bullet \text{ Soit } x_i = c_1 x^* + c_2, \quad x^* = \frac{1}{c_1} x_i - c_2$$

$$\leadsto \bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i = \frac{1}{n} \cdot \sum_{i=1}^n (c_1 x^* + c_2) = \left(\frac{c_1}{n} \cdot \sum_{i=1}^n x^*\right) + \frac{n}{n} \cdot c_2 = c_1 \cdot \bar{x}^* + c_2$$

$$\begin{aligned} \leadsto s^2 &= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \sum_{i=1}^n (c_1 x^* + c_2 - (c_1 \bar{x}^* + c_2))^2 = \frac{1}{n-1} \sum_{i=1}^n (c_1 x^* - c_1 \bar{x}^*)^2 = \\ &= \frac{c_1^2}{n-1} \cdot \sum_{i=1}^n (x^* - \bar{x}^*)^2 = c_1^2 \cdot \frac{1}{n-1} \cdot \sum_{i=1}^n (x^* - \bar{x}^*)^2 = c_1^2 \cdot (s^*)^2 \end{aligned}$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$x_i = c_1 x^* + c_2$$

Beh.: • **Thè.:**

$$1. \bar{x} = c_1 \cdot \bar{x}^* + c_2$$

$$2. s^2 = c_1^2 \cdot (s^*)^2$$

2.4.3 Berechnungen mittels der Häufigkeitsfunktion — Calculs au moyen de la fonction de fréquence

$$\text{Sei } \bullet \text{ Soit } \bar{x} = \bar{x} = \frac{1}{n} \cdot \sum_{i=1}^n x_i, \text{ Klassen } \bullet \text{ classes } \{[x]_1, \dots, [x]_m\}, \quad x_i \in [x]_j, \quad i \leq j.$$

Sei n_k die Häufigkeit in der Klasse k . Für die Klassen gilt dann:

• *Soit n_k la fréquence dans la classe k . Pour les classes vaut donc:*

$$\begin{aligned}\overline{[x]} &= \frac{1}{n} \cdot \sum_{k=1}^j n_k \cdot [x]_k, \quad h([x]_k) = \tilde{f}(x_k) = \frac{n_k}{n} \Rightarrow n_k = n \cdot \tilde{f}(x_k) \\ \Rightarrow \overline{[x]} &= \frac{1}{n} \cdot \sum_{k=1}^j n \cdot \tilde{f}(x_k) \cdot [x]_k = \sum_{k=1}^j \tilde{f}(x_k) \cdot [x]_k\end{aligned}$$

Bemerkung: • **Remarque:** $x_i \in [x]_k \Rightarrow x_i \approx x_k \wedge \overline{[x]} \approx \bar{x}$
 Fehler: Max. halbe Klassenbreite.
 • *Erreur: Au max. demi-largeur de classe.*

Satz: • **Théorème:** $\overline{[x]} = \sum_{k=1}^j \tilde{f}(x_k) \cdot [x]_k \approx \bar{x}$

Es gilt: • *Il vaut:*

$$\begin{aligned}s^2 &= \frac{1}{n-1} \sum_{i=1}^j (x_i - \bar{x})^2 \approx (\overline{[s]})^2 = \frac{1}{n-1} \sum_{i=k}^j ([x]_k - \overline{[x]})^2 \cdot n_k \\ &\approx \frac{1}{n-1} \sum_{i=k}^j ([x]_k - \overline{[x]})^2 \cdot n \cdot \tilde{f}([x]_k) = \frac{1}{n-1} \sum_{i=k}^j ([x]_k - \overline{[x]})^2 \cdot n \cdot \tilde{f}([x]_k)\end{aligned}$$

Auf Seite 21 haben wir berechnet:

• *A la page 21 nous avons calculé:*

$$s^2 = \frac{1}{n-1} \sum_{i=1}^j (x_i - \bar{x})^2 = \left(\frac{1}{n-1} \cdot \sum_{i=1}^n x_i^2 \right) - \left(\frac{n}{n-1} \cdot \bar{x}^2 \right) = \frac{1}{n-1} \cdot \left(\sum_{i=1}^n x_i^2 - \frac{1}{n} \cdot \left(\sum_{i=1}^n x_i \right)^2 \right)$$

Analog schliessen wir auch hier: • *Analogiquement nous déduisons aussi ici:*

Satz: • **Théorème:** Sei • *Soit* $s_K^2 := (\overline{[s]})^2$: $s^2 \approx s_K^2 = \frac{1}{n-1} \sum_{i=k}^j ([x]_k - \overline{[x]})^2 \cdot n \cdot \tilde{f}([x]_k)$

$$\begin{aligned}&= \frac{1}{n-1} \cdot \sum_{i=k}^j ([x]_k \cdot n \tilde{f}([x]_k) - \frac{1}{n} \cdot \left(\sum_{i=k}^j x_k n \tilde{f}([x]_k) \right)^2) \\ &= \frac{1}{n-1} \cdot \left(\sum_{i=k}^j ([x]_k) \right)^2 \cdot n \cdot \tilde{f}([x]_k) - n \cdot (\overline{[s]})^2\end{aligned}$$

2.4.4 Häufigkeitsverteilung und Massenverteilung — Distribution de fréquence et distribution de masse

Idee: • **Idée:**

Die Masse 1 sei längs einer Gerade verteilt. • *La masse 1 soit distribuée le long d'une droite.*

Werte: • <i>Valeurs:</i>	x_1	x_2		x_m
Häufigkeit: • <i>Fréquence:</i>	$h_1 = \tilde{f}(x_1)$	$h_2 = \tilde{f}(x_2)$		$h_m = \tilde{f}(x_m)$

Analogie: • *Analogiquement:*

Punkt: • <i>Point:</i>	x_1	x_2		x_m
Masse: • <i>Masse:</i>	h_1	h_2		h_m

$$\Rightarrow \tilde{F}(x_1) = \sum_{x_i \leq x} h_i = \sum \text{Massen links von } x \quad \bullet \quad \sum \text{masses à gauche de } x$$

Konsequenz: • Conséquence:

Der **Schwerpunkt** x_S entspricht dem **Mittelwert** $\bar{x}$:

• *Le centre de gravité* x_S *correspond à la valeur moyenne* $\bar{x}$: $x_S \hat{=} \bar{x} \cdot 1 = \sum x_i \cdot h_i$

Das **Trägheitsmoment** Φ entspricht bis auf einen Faktor der **Varianz** s^2 :

• *Le moment d'inertie* Φ *correspond, à l'exception d'un facteur, à la variance* s^2 : $\Phi \hat{=} \frac{n}{n-1} \cdot s^2$

2.4.5 Beispiel mit Mathematica — Exemple avec Mathematica

Packages laden • *Charger les "packages"* — Mathematica-Programm: • Programme en Mathematica:

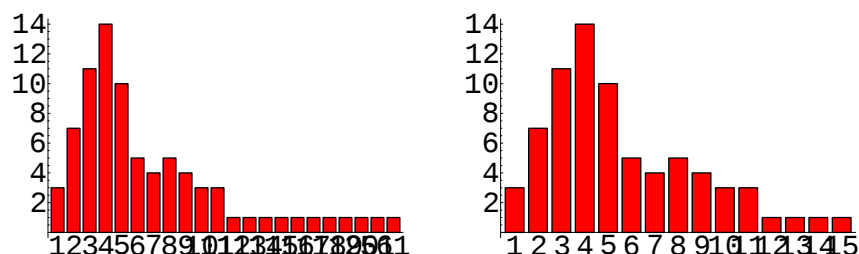
```
<< Statistics`DataManipulation`;
<< Statistics`DescriptiveStatistics`;
<< Graphics`Graphics`
```

Daten • *Données* — Mathematica-Programm: • Programme en Mathematica:

```
d = {0, 0, 0, 1, 1, 7, 7, 10, 14, 1, 1, 2, 2, 2, 2, 2, 1, 5, 5, 6, 3, 4, 2, 6,
      8, 9, 7, 13, 3, 4, 1, 2, 3, 5, 7, 5, 3, 3, 3, 3, 3, 3, 2, 4, 6, 8, 11,
      12, 4, 3, 3, 3, 3, 3, 4, 4, 4, 4, 4, 9, 7, 5, 6, 4, 2, 8, 2, 2, 1, 8, 9,
      10, 15, 16, 17, 18, 19, 10, 50, 60} + 1;
d1 = {0, 0, 0, 1, 1, 7, 7, 10, 6, 6, 1, 1, 2, 2, 2, 2, 2, 1, 5, 5, 6, 3, 4, 2,
       6, 8, 9, 7, 3, 4, 1, 2, 3, 5, 7, 5, 3, 3, 3, 3, 3, 3, 2, 4, 6, 8, 4, 3,
       3, 3, 3, 3, 4, 4, 4, 4, 4, 9, 7, 5, 6, 4, 2, 8, 2, 2, 1, 8, 9, 10, 10} + 1
```

Manipulation • *Manipulation* — Mathematica-Programm: • Programme en Mathematica:

```
f = Frequencies[d]; f1 = Frequencies[d1]; p1 = BarChart[f]; p2 = BarChart[f1];
Show[GraphicsArray[{p1, p2}]];
```



Mittelwerte • *Valeurs moyennes* — Mathematica-Programm: • Programme en Mathematica:

```
|"d", LocationReport[d] // N, "d1", LocationReport[d1] // N}|
```

Output: • Output:

```
|"d", {Mean -> 7.6125, HarmonicMean -> 4.13919, Median -> 5.},  
"d1", {Mean -> 5.16901, HarmonicMean -> 3.73698, Median -> 5.} }
```

Bemerkung: • **Remarque:** Der Median ist unempfindlich (robust) gegen Ausreisser!
 • *Le médian est indifférent (robuste) par rapport aux valeurs aberrantes!*

Manipulation • *Manipulation* — **Mathematica-Programm:** • **Programme en Mathematica:**

```
s = Sort[d];  
{s, "Median->", s[[Length[d]/2]], s[[Length[d]/2 + 1]], "Length->", Length[d]}
```

Output: • Output:

```
{ {1, 1, 1, 2, 2, 2, 2, 2, 2, 2, 3, 3, 3, 3, 3, 3, 3, 3, 3, 3, 4, 4, 4, 4,  
4, 4, 4, 4, 4, 4, 4, 4, 4, 4, 5, 5, 5, 5, 5, 5, 5, 5, 5, 5, 6, 6, 6, 6, 6,  
7, 7, 7, 7, 8, 8, 8, 8, 8, 9, 9, 9, 9, 10, 10, 10, 11, 11, 11, 12, 13,  
14, 15, 16, 17, 18, 19, 20, 51, 61},  
Median->, 5, 5, Length->, 80}
```

Streuungen • *Dispersions* — **Mathematica-Programm:** • **Programme en Mathematica:**

```
|"d", DispersionReport[d] // N, "d1", DispersionReport[d1] // N}|
```

Output: • Output:

```
{d, {Variance -> 80.215, StandardDeviation -> 8.95628, SampleRange -> 60.,  
MeanDeviation -> 4.92688, MedianDeviation -> 2., QuartileDeviation -> 3.},  
d1, {Variance -> 6.97103, StandardDeviation -> 2.64027, SampleRange -> 10.,  
MeanDeviation -> 2.15791, MedianDeviation -> 2., QuartileDeviation -> 2.}}
```

2.5 Auswertung: Beispiel — Exploitation: Exemple

2.5.1 Dateneingabe — Entrée de données

Wir verwenden *Mathematica*: • *Nous utilisons Mathematica*:

In:

```
Fichte={95.53,81.93,83.57,54.82,73.83,58.48,59.15,83.29};
```

```
Buche={113.14,156.61,169.96,142.20,162.85,196.08,125.03,88.47};
```

```
allData=Union[Fichte, Buche]
```

Out:

```
{54.82, 58.48, 59.15, 73.83, 81.93, 83.29, 83.57, 88.47, 95.53, 113.14,
 125.03, 142.2, 156.61, 162.85, 169.96, 196.08}
```

2.5.2 Kenngrößen — Caractéristiques

In:

```
<<Statistics'DescriptiveStatistics';
data=Fichte; {LocationReport[data], DispersionReport[data], ShapeReport[data]}
```

Out:

```
{{Mean -> 73.825, HarmonicMean -> 71.1504, Median -> 77.88},
 {Variance -> 219.052, StandardDeviation -> 14.8004, SampleRange -> 40.71,
  MeanDeviation -> 12.2563, MedianDeviation -> 11.67, QuartileDeviation -> 12.3075},
 {Skewness -> -0.0521399, QuartileSkewness -> -0.549055, KurtosisExcess -> -1.38182}}
```

In:

```
data=Buche; {LocationReport[data], DispersionReport[data], ShapeReport[data]}
```

Out:

```
{{Mean -> 144.293, HarmonicMean -> 136.328, Median -> 149.405},
 {Variance -> 1185.56, StandardDeviation -> 34.432, SampleRange -> 107.61,
  MeanDeviation -> 27.0825, MedianDeviation -> 22.465, QuartileDeviation -> 23.66},
 {Skewness -> -0.176882, QuartileSkewness -> -0.281488, KurtosisExcess -> -0.844303}}
```

In:

```
data=allData; {LocationReport[data], DispersionReport[data], ShapeReport[data] }
```

Out:

```
{{Mean -> 109.059, HarmonicMean -> 93.5017, Median -> 92.},
 {Variance -> 1979.66, StandardDeviation -> 44.4934, SampleRange -> 141.26,
  MeanDeviation -> 37.8073, MedianDeviation -> 32.94, QuartileDeviation -> 35.7625},
 {Skewness -> 0.521187, QuartileSkewness -> 0.605173, KurtosisExcess -> -0.996434}}
```

2.5.3 Darstellung mittels Kenngrößen: BoxWhiskerPlot — Représentation par caract.: BoxWhiskerPlot

Der **Box-Whisker-Plot** ist von unschätzbarem Wert um einen schnellen Überblick über einen numerischen Datensatz zu gewinnen. Er bekommt seine Form durch das Rechteck, das durch die Distanz zwischen zwei Quartilen um den Median herum gegeben wird. Üblicherweise sind es das 25% und das 75% Quartil. Normalerweise schliessen „whiskers“, d.h. Linien wie Backenbärte (Schnurrhaare) die Rechtecke ein, welche die Ausdehnung des Datensatzes entweder ganz oder ohne die Ausreisser markieren. **Ausreisser** sind Punkte jenseits einer Distanz von $3/2$ der zwischenquantilen Spanne weg vom Rechteck. **Grosse Ausreisser** sind Punkte jenseits von 3 mal diese **Spanne**.

Der **Quantilabstand (Quartilsdifferenz)** ist ein Mass der Streuung (Höhe des Rechtecks). Weiter ist der Median in der Box eingezeichnet, welcher durch seine Lage innerhalb der Box einen Eindruck von der Schiefe der den Daten zugrunde liegenden Verteilung vermittelt. Diese Quartile (Quantile) sowie der

Median sind robust gegen Ausreisser: Sie ändern nicht enorm, wenn man die Ausreisser weglässt, im Gegensatz zum arithmetischen Mittelwert und damit der Standardabweichung.

Wie schon der **Median** ein **robustes Lagemass** war, so ist somit auch die **Quartilsdifferenz** ein **robustes Streumass gegen Ausreisser**.

• *Le Box–Whisker–Plot est d’une valeur inestimable pour gagner rapidement une vue générale sur un paquet de données numériques. Ce paquet reçoit sa forme par un rectangle, qui est donné autour du médian par la distance entre les deux quartiles. Normalement ce sont les quartiles de 25 % et de 75 %. Normalement ce sont des “whiskers”, c.-à.-d. des lignes comme les poils d’une moustache qui enferment un rectangle, qui marque l’élargissement du paquet des données soit entièrement, soit sans les valeurs aberrantes. Les points aberrants sont des points au-delà d’une distance de $3/2$ de la distance entre les quartiles mentionnées, mesurée depuis le rectangle. Puis les **grandes valeurs aberrantes** sont marquées par des points au-delà de 3 fois cette marge.*

*La **distance des quartiles (différence des quartiles)** est une mesure de la dispersion ou variance (hauteur du rectangle). En outre le médian est inscrit dans le rectangle et donne par sa situation dans le rectangle une impression du biais de la distribution des données. Ces quartiles ainsi que le médian résistent aux valeurs aberrantes: Elles ne changent pas énormément si on omet les valeurs aberrantes, contrairement à la moyenne arithmétique et à l’écart standard. Comme déjà le **médian** était une mesure de position **robuste**, ainsi par conséquent aussi la différence des quartiles est une mesure de l’écart robuste contre les valeurs aberrantes.*

Auszug aus dem Help–File von *Mathematica*: • *Extrait du fichier Help de Mathematica:*

„The box-and-whisker plot is invaluable for gaining a quick overview of the extent of a numeric data set. It takes the form of a box that spans the distance between two quantiles surrounding the median, typically the 25 % quantile to the 75 % quantile. Commonly, „whiskers“, lines that extend to span either the full data set or the data set excluding outliers, are added. Outliers are defined as points beyond $3/2$ the interquantile range from the edge of the box; far outliers are points beyond three times the interquantile range.“

Bemerkung: • **Remarque:** Bei den **Quartilen** benutzt man die folgende Definition:

• *Quant aux quartiles on utilise la définition suivante:*

Definition: • **Définition:** **Quartil q_j :** • **Quartile q_j :**

Sei • *Soit $n = 4k + r$, $k, r \in \mathbb{N}$, x_k = Messwert Nummer k in der Rangliste der Messwerte • *valeur de mesure numéro k dans le classement des valeurs indiquées $\leadsto$**

$$r = 0 \Rightarrow q_{0.25} = 0.25 x_k + 0.75 x_{k+1}$$

$$r = 1 \Rightarrow q_{0.25} = x_{k+1}$$

$$r = 2 \Rightarrow q_{0.25} = 0.75 x_{k+1} + 0.25 x_{k+2}$$

$$r = 3 \Rightarrow q_{0.25} = 0.5 x_{k+1} + 0.5 x_{k+2}$$

$$r = \dots \Rightarrow q_{0.25} = \tilde{x}$$

$$r = 0 \Rightarrow q_{0.75} = 0.75 x_{3k} + 0.25 x_{3k+1}$$

$$r = 1 \Rightarrow q_{0.75} = x_{3k+1}$$

$$r = 2 \Rightarrow q_{0.75} = 0.25 x_{3k+1} + 0.75 x_{3k+2}$$

$$r = 3 \Rightarrow q_{0.75} = 0.5 x_{3k+2} + 0.5 x_{3k+3}$$

$Q_1 = q_{0.25} = \mathbf{1. Quartil}$, $Q_2 = q_{0.5} = \mathbf{2. Quartil} = \text{Median}$, $Q_3 = q_{0.75} = \mathbf{3. Quartil}$, entsprechend für **Quintile** u.s.w.

Das erste Quartil ist also der Wert bei etwa einem Viertel der Messwerte, das zweite Quartil der Wert bei etwa der Hälfte der Messwerte und das dritte Quartil der Wert bei etwa dreien Vierteln der Messerte.

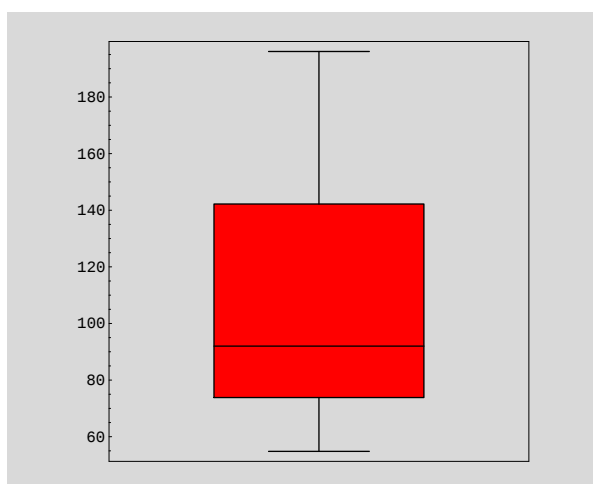
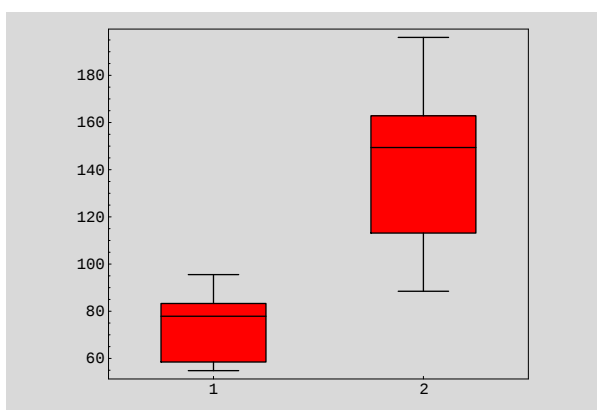
• $Q_1 = q_{0.25} = \mathbf{1er quartile}$, $Q_2 = q_{0.5} = \mathbf{2ème quartile} = \textit{médian}$, $Q_3 = q_{0.75} = \mathbf{3ème quartile}$, *correspondant pour les quintiles etc.*

La première quartile est donc la valeur environ à un quart des valeurs indiquées, la deuxième quartile est la valeur environ à la moitié des valeurs indiquées et la troisième quartile est la valeur à trois quart des valeurs indiquées.

Vgl. auch: • Voir aussi: <http://de.wikipedia.org/wiki/Boxplot>

In:

```
<<Statistics`StatisticsPlots`
BoxWhiskerPlot[Transpose[{Fichte,Buche}]];      BoxWhiskerPlot[Transpose[{allData}]];
```



2.5.4 Andere statistische Plots — D'autres représentation statistiques

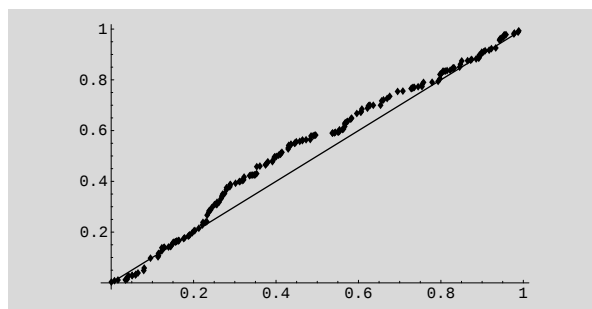
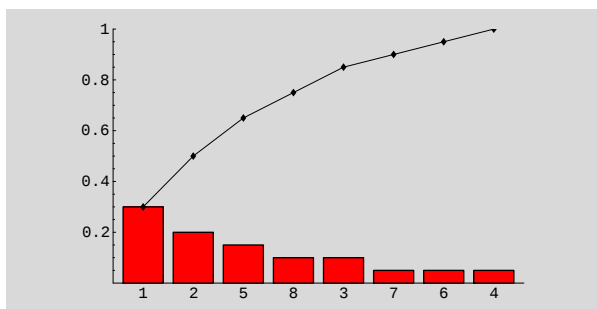
Andere statistische Plots: Vgl. Erklärungen in *Mathematica*.

• D'autres plots statistiques: Voir explications de *Mathematica*.

In:

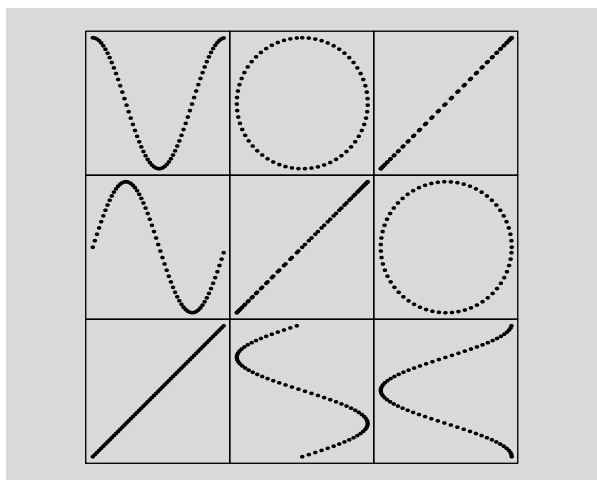
```
data={1,1,1,2,2,3,2,4,1,5,5,6,3,1,1,7,2,8,8,5};
```

```
ParetoPlot[data];      QuantilePlot[Table[Random[],{200}],Table[Random[],{200}]];
```



In:

```
data = Table[{x, Sin[x], Cos[x]}, {x, 0, 2 Pi, 0.1}]; PairwiseScatterPlot[data];
```



In:

```
StemLeafPlot[data];
```

Siehe Help in *Mathematica*. • Voir Help dans *Mathematica*.

2.6 Weitere Kenngrößen — D'autres caractéristiques

2.6.1 Diverse Mittelwerte einer Verteilung — Certaines valeurs moyennes d'une distribution

Definitionen und Folgerungen:

• Définitions et conclusions:

1. **Arithmetisches Mittel:**

• Moyenne arithmétique:

$$\bar{x}_a = \frac{1}{n} \sum_{j=1}^n x_j = \frac{x_1 + x_2 + \dots + x_n}{n}$$

2. **Gewichtetes arithmetisches Mittel, Gewicht w_i zu x_i :**

• Moyenne arithmétique pondérée, poids w_i à x_i :

$$\bar{x}_{a, gew} = \frac{1}{\sum_{j=1}^n w_j} \cdot \sum_{j=1}^n w_j \cdot x_j = \frac{w_1 \cdot x_1 + w_2 \cdot x_2 + \dots + w_n \cdot x_n}{\sum_{j=1}^n w_j}$$

3. **Geometrisches Mittel:**

• Moyenne géométrique:

$$\bar{x}_g = \left(\prod_{j=1}^n x_j \right)^{\frac{1}{n}} = \sqrt[n]{\prod_{j=1}^n x_j} = \sqrt[n]{x_1 \cdot x_2 \cdot \dots \cdot x_n}, \quad \ln(\bar{x}_g) = \frac{1}{n} \sum_{j=1}^n \ln(x_j) = \overline{\ln(x)}_a$$

4. **Gewichtetes geometrisches Mittel:**

• Moyenne géométrique pondérée:

$$\bar{x}_{g, gew} = \left(\prod_{j=1}^n x_j^{w_j} \right)^{\frac{1}{w}} = \sqrt[w]{\prod_{j=1}^n x_j^{w_j}}, \quad w = \sum_{j=1}^n w_j$$

5. Harmonisches Mittel:

- Moyenne harmonique:

$$\bar{x}_h = \frac{n}{\sum_{j=1}^n \frac{1}{x_j}}, \quad \sum_{j=1}^n \frac{1}{\bar{x}_h} = \frac{n}{\bar{x}_h} = \sum_{j=1}^n \frac{1}{x_j}$$

Für nur zwei Werte x_1 und x_2 gilt:

- Pour seulement deux valeurs x_1 et x_2 il vaut:

$$\bar{x}_h = \frac{\bar{x}_g^2}{\bar{x}_a}$$

6. Gewichtetes harmonisches Mittel:

- Moyenne harmonique pondérée:

$$\bar{x}_{h, gew} = \frac{\sum_{j=1}^n w_j}{\sum_{j=1}^n \frac{w_j}{x_j}}$$

7. Logarithmisches Mittel von zwei positiven Werten:

- Moyenne logarithmique de deux valeurs positives:

$$\bar{x}_{\ln, x_1, x_2} = \frac{x_1 - x_2}{\ln(x_1) - \ln(x_2)}$$

Dann gilt: • Alors il vaut:

$$\bar{x}_g < \bar{x}_{\ln, x_1, x_2} < \bar{x}_a$$

8. Potenzmittel mit positiven Werten:

- Moyenne de puissances avec des valeurs positives:

$$\bar{x}_{p, k} = \sqrt[k]{\frac{1}{n} \sum_{j=1}^n x_j^k} = \left(\frac{1}{n} \sum_{j=1}^n x_j^k \right)^{\frac{1}{k}}$$

Bemerkung: • Remarque:

Für $k = 1$ hat man das arithmetische Mittel, für $k = 2$ das quadratische Mittel $\bar{x}_q$, für $k = -1$ das harmonische Mittel u.s.w.

Weiter gilt mit dem r -ten statistischen Moment (siehe nächster Abschnitt):

- Pour $k = 1$ on a la moyenne arithmétique, pour $k = 2$ on a la moyenne quadratique $\bar{x}_q$, pour $k = -1$ on a la moyenne harmonique etc.

En outre avec le r -ième moment statistique (voir section suivante):

$$m_r = (\bar{x}_{p, r})^r \text{ sowie } \bullet \text{ ainsi que } u < v \Rightarrow \bar{x}_{p, u} \leq \bar{x}_{p, v}$$

$$x_{\min} \leq \bar{x}_h \leq \bar{x}_g \leq \bar{x}_a \leq \bar{x}_q \leq x_{\max}$$

9. Quasi-arithmetisches Mittel: Sei f streng monoton und stetig, $w_j \in [0, 1]$, $\sum_{j=0}^n w_j = 1$

- Moyenne quasiment arithmétique: f soit continue de façon strictement monotone et continue, $w_j \in [0, 1]$, $\sum_{j=0}^n w_j = 1$

$$\bar{x}_f = f^{-1} \left(\sum_{j=1}^n w_j \cdot x_j \right)$$

10. **Winsorisierte oder gestutzte Mittelwerte:** Wenn die Daten durch Ausreisser (einige wenige viel zu hohe oder viel zu niedrige Werte) verunreinigt sind und man diese Werte nicht berücksichtigen will, so „winsorisiert“ man. D.h. man sortiert die Beobachtungswerte nach aufsteigendem Rang oder Grösse. Dann streicht man am Anfang und am Ende der so entstehenden Rangliste eine gleiche Anzahl von Werten, die Ausreisser also, weg. Darauf berechnet man von den verbleibenden Werten den Mittelwert.

Beispiel: Ein 20 % winsorisierter Mittelwert wird erzeugt, indem man 10 % der Anzahl aller Werte am unteren Ende und 10 % dieser Werte am oberen Ende der Rangliste streicht.

• **Des moyennes "winsorisées":** *Si les données sont contaminées par des valeurs aberrantes (quelques valeurs beaucoup trop hautes ou beaucoup trop basses) et si on ne veut pas considérer ces valeurs comme données valables, on "winsorise" les données. C.-à.-d. on assortit les valeurs d'observation d'après le rang ascendant ou descendant. Alors on trace un même nombre de valeurs au début et à la fin du classement si obtenu, de façon que les valeurs aberrantes disparaissent. Après on calcule la moyenne sur des valeurs restantes. Exemple: On produit une moyenne "winsorisée" du niveau de 20 % en traçant 10 % des nombres de toutes les valeurs à la fin inférieure et 10 % des valeurs à la fin supérieure du classement.*

11. Weitere interessante Mittelwerte, von denen man jedoch einzelne selten in der deskriptiven Statistik trifft, sind das ***a*-Mittel**, verschiedene Ausprägungen von **gleitenden Mittelwerten** oder das **arithmetisch-geometrische Mittel**. Bezüglich dieser Mittelwerte sei auf die einschlägige Fachliteratur verwiesen, da sie den Rahmen dieses Kurses sprengen.

• *D'autres moyennes intéressantes, cependant qu'on rencontre rarement dans la statistique descriptive, sont la moyenne "a" et des manifestations distinctes de moyennes glissantes ou de moyennes arithmétiques-géométriques. Concernant ces moyennes, le lecteur est prié de consulter la littérature spécialisée parce qu'elles excèdent le cadre de ce cours.*

Hinweis:

- *Indication:*

Mit dem Höhensatz, dem Kathetensatz und

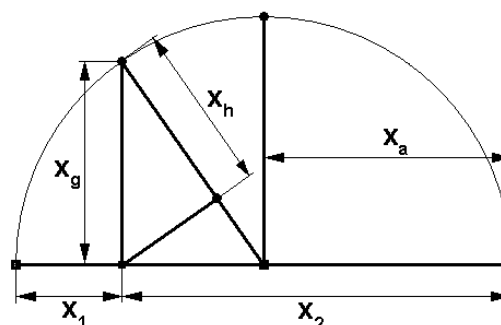
- *A l'aide du théorème de la hauteur et*

$$\bar{x}_h = \frac{\bar{x}_g^2}{\bar{x}_a}$$

sieht man aus nebenstehender Figur:

- *On voit dans la figure ci-contre:*

$$\bar{x}_h \leq \bar{x}_g \leq \bar{x}_a$$



2.6.2 Momente einer Verteilung — Moments d'une distribution

Seien $x_1, x_2, \dots, x_n$ Messwerte, die mehrmals vorkommen können. Wir nehmen an, dass, wenn wir nur jeweils ein Exemplar jedes vorkommenden Messwerts auflisten, wir die Liste $\{x_1, x_2, \dots, x_k\}$ bekommen. Seien $f_1 = n_1, f_2 = n_2, \dots, f_k = n_k$ die zur Liste gehörigen Frequenzen oder absolute Häufigkeiten. Dann gilt für die totale Anzahl der Messwerte $n = n_1 + n_1 + \dots + n_k = f_1 + f_1 + \dots + f_k$.

• *Soyent $x_1, x_2, \dots, x_n$ les valeurs indiquées, qu'on rencontre plusieurs fois. Nous supposons que, si nous faisons une liste avec chaque fois seulement un exemplaire de chaque valeur rencontrée, nous obtenons la liste $\{x_1, x_2, \dots, x_k\}$. Soyent $f_1 = n_1, f_2 = n_2, \dots, f_k = n_k$ les fréquences ou fréquences absolues dues à la liste. Alors il vaut pour le nombre total des valeurs rencontrées $n = n_1 + n_1 + \dots + n_k = f_1 + f_1 + \dots + f_k$.*

Analog zum **Mittelwert** definieren wir jetzt das **Moment *r*-ter Ordnung** $m_{r,a}$ der Messreihe in Bezug auf einen Ausgangspunkt oder Ausgangswert a . Wenn $a = \bar{x}$ ist, so schreiben wir kurz m_r und

sprechen vom **Moment r -ter Ordnung**, ohne den Ausgangspunkt zu nennen.

- *Analogiquement à la moyenne on définit maintenant $m_{r,a}$ comme moment d'ordre r de la série de mesure par référence à un point de départ ou valeur initiale "a". S'il vaut $a = \bar{x}$, nous écrivons brièvement m_r et nous parlons ici du moment d'ordre r sans mentionner le point de départ.*

Definition: • **Définition:**

$$m_{r,a} = \frac{1}{n} \sum_{j=1}^k n_j (x_j - a)^r = \frac{n_1 (x_1 - a)^r + n_2 (x_2 - a)^r + \dots + n_k (x_k - a)^r}{n} = \overline{(x - a)^r}$$

Folgerungen:

• **Conclusions:**

Wenn $a = \bar{x}$ gilt, so ergibt sich:

- *S'il vaut $a = \bar{x}$, on obtient:*

$$m_r = \frac{1}{n} \sum_{j=1}^k n_j (x_j - \bar{x})^r = \frac{n_1 (x_1 - \bar{x})^r + n_2 (x_2 - \bar{x})^r + \dots + n_k (x_k - \bar{x})^r}{n} = \overline{(x - \bar{x})^r}$$

Wenn dann noch $f_1 = n_1 = f_2 = n_2 = \dots = f_k = n_k = 1$ gilt, so ergibt sich:

- *S'il vaut en plus $f_1 = n_1 = f_2 = n_2 = \dots = f_k = n_k = 1$, on obtient:*

$$m_r = \frac{1}{n} \sum_{j=1}^n (x_j - \bar{x})^r = \frac{(x_1 - \bar{x})^r + (x_2 - \bar{x})^r + \dots + (x_n - \bar{x})^r}{n} = \overline{(x - \bar{x})^r}$$

Wenn $a = 0$ gilt, so ergibt sich:

- *S'il vaut $a = 0$, on obtient:*

$$m_{r,0} = \frac{1}{n} \sum_{j=1}^k n_j x_j^r = \frac{n_1 x_1^r + n_2 x_2^r + \dots + n_k x_k^r}{n} = \overline{x^r}$$

Daher gilt für den Mittelwert $\bar{x}$ und die Standardabweichung s :

- *Par conséquent il vaut pour la moyenne $\bar{x}$ et l'écart standard s :*

$$m_{1,0} = \bar{x}, \quad s^2 = \frac{n}{n-1} \cdot m_{2,0}$$

Weiter gelten die Zusammenhänge:

- *En outre on obtient les relations suivantes:*

1. $m_{1,a} = \bar{x} - a$
2. $m_2 = m_{2,a} - m_{1,a}^2$
3. $m_3 = m_{3,a} - 3 m_{1,a} m_{2,a} + 2 m_{1,a}^3$
4. $m_4 = m_{4,a} - 4 m_{1,a} m_{3,a} + 6 m_{1,a}^2 m_{2,a} - 3 m_{1,a}^4$
5. ...

Bemerkung: • **Remarque:**

In der Literatur findet man eine grosse Fülle von Behandlungen der Momente. Zu erwähnen sind dabei „Charlier's Probe“ und „Sheppard's Korrektur“ sowie die „dimensionslose Form der Momentendefinition“ $M_r = \frac{m_r}{s^r}$. (Vgl. z.B. Murray R. Spiegel, Statistik, McGraw-Hill, Schaum-Lectures.)

- *Dans la littérature, on trouve une grande abondance de traitements des moments. Nous mentionnons "l'épreuve de Charlier" et "la correction de Sheppard" ainsi que la forme sans dimension de la définition des moments $M_r = \frac{m_r}{s^r}$. (Voir p.ex. Murray R. Spiegel, Statistique, McGraw-Hill, Schaum-Lectures.)*

2.6.3 Die Schiefe einer Verteilung — Le biais d'une distribution

Die **Schiefe** soll die Stärke der Einseitigkeit oder Neigung einer Häufigkeitsfunktion oder einer Verteilung beschreiben. Schiefe ist also ein Mass für die Asymmetrie. In manchen Fachgebieten definiert man ein solches Maß nach den praktischen Bedürfnissen. Ein offensichtlicher physikalischer Zugang zur Schiefe ist nicht unmittelbar ersichtlich. Die Schiefe ist in der Regel keine Messgröße. Die Schiefe gibt an, ob (und wie stark) eine Verteilung nach rechts (positiv) oder nach links (negativ) geneigt ist. Es ist leicht nachvollziehbar, dass die nachfolgenden Definitionen für die deskriptive Statistik Sinnvolles leisten:

• *Le biais devrait décrire la mesure du caractère unilatéral ou la déclivité d'une fonction de fréquence. Le biais est ainsi une mesure pour l'asymétrie. Dans certains domaines spécialisés, on définit une telle mesure d'après les besoins pratiques. Un accès physique évident au biais n'est pas directement visible. Le biais n'est normalement pas une grandeur à mesurer. Le biais indique, si (et avec quelle force) une distribution est penchée à droite (positivement) ou à gauche (négativement). Il est facilement compréhensible que les définitions suivantes sont raisonnables pour la statistique descriptive:*

Definition: • **Définition:**

Erster Pearsonscher Schiefekoeffizient:

• **Premier coefficient de biais de Pearson:**

$$P_{S,1} := \frac{\text{Mittelwert} - \text{Modus}}{\text{Standardabweichung}} = \frac{\bar{x} - x_{\text{modus}}}{s}$$

$$\bullet P_{S,1} := \frac{\text{valeur moyenne} - \text{modus}}{\text{écart-type}} = \frac{\bar{x} - x_{\text{modus}}}{s}$$

Zweiter Pearsonscher Schiefekoeffizient:

• **Deuxième coefficient de biais de Pearson:**

$$P_{S,1} := \frac{3 \cdot (\text{Mittelwert} - \text{Median})}{\text{Standardabweichung}} = 3 \cdot \frac{\bar{x} - x_{\text{Median}}}{s}$$

$$\bullet P_{S,1} := \frac{3 \cdot (\text{valeur moyenne} - \text{médian})}{s} = 3 \cdot \frac{\bar{x} - x_{\text{median}}}{s}$$

Ein Momentenkoeffizient als Schiefenmass:

• **Un coefficient de moment comme mesure de biais:**

$$a_3 := \frac{m_3}{s^3} = \frac{m_3}{m_2^{\frac{3}{2}}}$$

Hinweis: Der Median ist einfacher greifbar als ein Modus, vor allem wenn es sich um multimodale Verteilungen handelt.

• *Indication: Le médian est disponible de façon plus simple que le mode, en particulier s'il s'agit d'une distribution multi-modale.*

2.6.4 Kurtosis und Exzess — Kurtosis et excès

Die **Kurtosis** und der **Exzess** beschreiben die Wölbung oder Spitzigkeit einer statistischen Verteilung. Schmale aber hohe Verteilungen nennt man „leptokurtisch“. Flache, niedere Verteilungen heissen „platykurtisch“ und mittlere Verteilungen, etwa wie eine schöne Normalverteilung, heissen „mesokurtisch“.

Der mit der Kurtosis verknüpfte Exzess beschreibt die Abweichung im Verlauf der Verteilung zum Verlauf einer Normalverteilung.

• *La "kurtosis" et "l'excès" décrivent le bombement ou le caractère pointu d'une distribution statistique. On appelle une distribution "leptokurtique" si elle est haute et mince. On appelle une distribution "platykurtique" si elle est inférieure et plate. Des distributions à peu près normales s'appellent*

”mesokurtiques”. L’excès lié avec la kurtosis décrit l’écart sur le tracé de la distribution par rapport au tracé d’une distribution normale.

Definition: • **Définition:** **Momenten–Kurtosis:** • **Kurtosis des moments:**

$$\beta_2 := \frac{m_4}{s^4} = \frac{m_4}{m_2^2}$$

Quartil–Dezil–Kurtosis: • **Kurtosis quartile–décile:**

Analog zu den Quartilen Q_1, Q_2, Q_3 definiert man auch die Zentile $P_1, P_2, \dots, P_{99}$. Sei Q der Quantilabstand: $Q = \frac{1}{2}(Q_3 - Q_1)$. Dann definiert man die Quartil–Dezil–Kurtosis:

• *Analogiquement aux quartiles Q_1, Q_2, Q_3 on définit aussi les centiles $P_1, P_2, \dots, P_{99}$. Soit Q la distance des quartiles: $Q = \frac{1}{2}(Q_3 - Q_1)$. Alors on définit la kurtosis quartile–décile:*

$$\kappa := \frac{Q}{P_{90} - P_{10}}$$

Exzess: • **Excès:**

$$\gamma_2 = \beta_2 - 3$$

2.6.5 Sinn und Gefahr von Kenngrößen — Caractéristiques: Sens propre et danger

Ein Beispiel mag das Problem der Sinnhaftigkeit von Kenngrößen illustrieren. In der Praxis trifft man oft die Situation, dass Messwerte vorhanden sind, die nun auf allerlei Art höchst kopflos verrechnet werden. Angenommen, es gehe um den Vergleich zwischen dem Wohlbefinden von zwei Personen, Hans und Franz. Hans liegt auf dem Bett. In seinem Zimmer misst man die Temperatur von $30^\circ C$. Die Temperatur stimmt. Für das Wohlbefinden von Hans gibt es bezüglich der Zimmertemperatur kein Problem. Der Mittelwert der Temperatur seiner beiden Hände ist $30^\circ C$.

In einem andern Haus in einem andern Zimmer wird Franz von zwei Gangstern festgehalten. Er wird gefoltert. Eine Hand hat man ihm in kochendes Wasser gedrückt. Wasser kocht bei $100^\circ C$. An der Handoberfläche, wo die Nervenzellen zur Temperabschätzung sitzen, herrscht somit $100^\circ C$. Die andere Hand hat man ihm in einem Eisklotz bei $-40^\circ C$ eingefroren. An jener Handoberfläche herrscht somit $-40^\circ C$. Der Durchschnitt, der Mittelwert also, ist $\frac{1}{2} \cdot (-40^\circ C + 100^\circ C) = 30^\circ C$. Die so ermittelte Durchschnittstemperatur bei Franz ist also gleich der bei Hans. Wer nichts von der Sache weiss, nur den jeweiligen Mittelwert von $30^\circ C$ kennt und sich auf diese eine statistische Aussage verlässt, muss also zum Schlusse kommen, dass die beiden, Hans und Franz, bezüglich Temperatur–Wohlbefinden sich nicht unterscheiden. Krass! Man bedenke, dass täglich mit statistischen Kenngrößen Politik und Werbung gemacht wird, ohne dass dabei jemals weitere Worte über die Datenbestände und ihre Umstände fallen! — !! ? !!

• *Un exemple peut illustrer le problème du sens des mesures statistiques. En pratique, on rencontre souvent la situation que les valeurs indiquées sont utilisées de manière irréfléchie dans des calculs de toutes sortes. Supposons qu’il s’agisse de traiter le bien-être de deux personnes, Jean et Francis. Jean est couché sur le lit. Dans sa chambre, on mesure la température de $30^\circ C$. La température est en ordre. Pour le bien-être de Jean, il n’y a pas de problème concernant la température de la chambre. La valeur moyenne de la température de ses deux mains est de $30^\circ C$.*

Dans une autre maison dans une autre chambre, Francis est retenu par deux gangsters. Il est torturé.

On lui a mis de façon serrée une main dans de l'eau bouillante. L'eau cuit à $100^{\circ}C$. À la surface de la main, où les neurones sont placés pour la sensation de la température, on constate par conséquent $100^{\circ}C$. On lui a gelé l'autre main dans un bloc de glace à $-40^{\circ}C$. À cette surface de la main on constate par conséquent $-40^{\circ}C$. La moyenne est $\frac{1}{2} \cdot (-40^{\circ}C + 100^{\circ}C) = 30^{\circ}C$. La température moyenne est donc la même chez Francis est chez Jean. Celui qui ne connaît pas la situation des deux et qui connaît seulement la moyenne de $30^{\circ}C$ et qui se fie à cette déclaration statistique, doit croire que les deux, Jean et Francis, concernant leur bien-être par rapport à la température, ne se distinguent pas. Ahurissant! On pense que pratiquement chaque jour on fait de la politique et de la publicité avec des données statistiques, sans que jamais un mot ne soit dit sur l'ensemble des données et leurs circonstances!

Kapitel 3

Kombinatorik — Analyse combinatoire

3.1 Einleitung — Introduction

3.1.1 Problemstellung — Problème

Im Stoffgebiet *Kombinatorik* behandeln wir die 6 Typen der klassischen *Anzahlprobleme*. Dabei geht es um folgende Kategorie von Fragestellungen: Wir fragen nach der *Anzahl* der Möglichkeiten, aus einer endlichen Menge M nach einer gewissen Vorschrift Elemente *auszuwählen*, diese ev. *anzuordnen* oder die Menge *in Klassen einzuteilen*. Dabei können sich Elemente *wiederholen* – oder nicht. Da das Resultat y jeweils eine natürliche Zahl ist, reden wir auch von **Anzahlfunktionen** $M \mapsto y$.

• *Dans l'analyse combinatoire, nous traitons les 6 types de problèmes des nombres cardinaux classiques. Il s'agit des catégories suivantes de questions: Nous demandons le nombre des possibilités de pouvoir choisir des éléments dans un ensemble fini M d'après une prescription donnée, et éventuellement aussi de les ordonner ou de diviser l'ensemble en classes. Dans certains cas les éléments peuvent être répétés — ou bien non répétés. Comme le résultat y est chaque fois un nombre naturel, nous parlons de **fonctions dans les nombres cartinaux** $M \mapsto y$.*

3.1.2 Fakultäten — Factorielles

In der Kombinatorik spielt der Begriff *Fakultät* eine grosse Rolle. Man definiert die *Fakultäten induktiv*¹ durch die folgende *Rekursion*:

• *Dans l'analyse combinatoire, la notion des factorielles joue un grand rôle. On définit les factorielles*² *par la relation de récurrence suivante:*

Definition • **Définition 3.1 (Fakultät: • Factorielle:)**

$$\begin{array}{ll} f(0) = 0! := 1 & (\text{Verankerung}) \bullet (\text{Ancrage}) \\ f(n) = n! := n \cdot (n-1)! & (\text{Vererbung}) \bullet (\text{Hérédité}) \end{array}$$

Bemerkungen: • *Remarques:*

¹Nach dem Schema der vollständigen Induktion, vgl. Thema *natürliche Zahlen*, *Induktionsaxiom* (eines der Axiome aus dem System von Peano).

²D'après le schéma de l'induction complète, voir le sujet des nombres naturels, axiome d'induction, un des axiomes du système de Peano.

1. Es gilt dann: • *Il vaut donc:* $n! = n \cdot (n-1) \cdot (n-2) \cdot \dots \cdot 3 \cdot 2 \cdot 1 = \prod_{k=1}^n k$. (Siehe • *Voir*⁽³⁾.)
Daraus ergibt sich: • *Il vaut donc:* $1! = 1$, $2! = 2$, $3! = 6$, $4! = 24$ etc..
2. Der Begriff *Rekursion* hat sich heute in der *Informatik* sehr stark verbreitet. Man versteht dort darunter *die Definition einer Funktion oder eines Verfahrens durch sich selbst* (vgl. dazu z.B. Bibl. A6, (Claus, Schwill)). Man darf den Begriff *Rekursion* in diesem hier verwendeten einfachen Sinne jedoch nicht verwechseln mit den in der höheren Mathematik gebräuchlichen, etwas schwierigen Begriffen *allgemeine Rekursion*, *primitive Rekursion*, *rekursive Funktion* (in der Zahlentheorie), *rekursive Relation* (in verschiedenem Sinne in Logik und Mengenlehre). Vgl. dazu Fachlexikon a b c (Bibl.: A1), Iyanaga, Kawada (Bibl. A8) und Meschkowski, Bibl. A11.
• *La notion de récurrence est très répandue aujourd'hui dans l'informatique. On entend par cette notion la définition d'une fonction ou une méthode par elle-même (voir aussi par exemple Bibl. A6 (Claus, Schwill)). Mais il ne faut pas confondre la notion de la récurrence qui cependant est utilisée ici dans un sens simple avec les notions récurrence commune, récurrence primitive, fonction de récurrence (dans la théorie des nombres) et relation de récurrence (dans des sens différents dans la logique et la théorie des ensembles) qui sont un peu difficiles et usuelle dans les mathématiques. Voir aussi Fachlexikon a b c (Bibl.: A1), Iyanaga, Kawada (Bibl. A8) et Meschkowski, Bibl. A11.*

Wir halten fest: • *Nous retenons:*

Definition • **Définition 3.2 (Rekursion (Informatik) • Récurrence (informatique))**
Unter Rekursion verstehen wir hier die Definition einer Funktion oder eines Verfahrens durch sich selbst.
• *Sous récurrence nous entendons la définition d'une fonction ou d'une méthode par elle-même.*

Beispiel: • **Exemple:** Aus der Definition von $n!$ ergibt sich: $f(n) = n \cdot f(n-1)$. Die Funktion an der Stelle n wird also durch die Funktion (also durch sich selbst) an der Stelle $n-1$ definiert. • *De la définition de $n!$ on conclut: $f(n) = n \cdot f(n-1)$. La fonction à la place n est donc définie par la fonction à la place $n-1$ (ainsi par elle-même).*

Die Werte $f(n) = n!$ werden sehr rasch sehr gross. z.B. ist: • *Les valeurs $f(n) = n!$ augmentent très vite. Par exemple il vaut:*

$$40! = 815915283247897734345611269596115894272000000000 \approx 8.15915 \cdot 10^{47}.$$

Ein einfaches Programm auf einem Rechner kann daher schnell Probleme machen. Hier ist eine Formel von *Stirling* hilfreich (ohne Beweis): • *Un programme simple sur un ordinateur peut donc très vite causer des problèmes. Voici une formule de Stirling très utile (sans la preuve):*

Satz • **Théorème 3.1 (Formel von Stirling: • Formule de Stirling:)**

$$n! \approx \sqrt{2\pi n} \left(\frac{n}{e}\right)^n$$

e ist hier die Eulersche Zahl: • *Ici e est le nombre de Euler: $e \approx 2.71828182845904523536028747135$ (auf 30 Stellen • à 30 places)*

3.2 Anordnungsprobleme — Problèmes d'arrangement

3.2.1 Permutationen ohne Wiederholung — Permutations sans répétition

Paradigma — **Paradigme (Beispiel eines praktischen Problems)** • **(Exemple d'un problème pratique)**⁴

Problem • **Problème 3.1 (Sitzmöglichkeiten: • Possibilités de s'asseoir:)**

³ $\prod$ steht für „Produkt“. • $\prod$ signifie „produit“.

⁴Ein Paradigma ist ein Lehrbeispiel • *Un paradigme est un exemple démonstratif*

- Situation:** • **Situation:** *In einem Klassenzimmer befindet sich nichts ausser 26 nummerierten Stühlen. Die Nummern gehen von 1 bis 26. Pulte, Bänke und Tische hat man nach draussen gebracht. Vor der Tür warten 26 Studenten. Zur besseren Unterscheidbarkeit und Benennung erhält auch jeder Student eine verschiedene Nummer von 1 bis 26, mit der er aufgerufen wird.*
- *Dans une salle de classe il ne se trouve rien à l'exception de 26 chaises numérotées. Les numéros vont de 1 à 26. On vient d'enlever pupitres, bancs et tables. 26 étudiants attendent devant la porte. Pour pouvoir mieux distinguer les étudiants et pour mieux pouvoir les appeler, chaque étudiant reçoit un numéro différent de 1 à 26 avec lequel il est donc appelé.*
- Frage:** • **Question:** *Auf wieviele Arten kann man die 26 Studenten auf die 26 Stühle setzen, d.h. wieviele Sitzordnungen gibt es?*
- *De combien de manières différentes est-ce qu'on peut mettre les 26 étudiants sur les 26 chaises, c.-à.-d. combien est-ce que de répartitions des places existent?*

Lösung: • **Solution:**

- Der Student Nr. 1 kommt herein. Er findet 26 freie Stühle vor. Somit hat er für sich 26 Sitzmöglichkeiten.
- *L'étudiant no. 1 entre. Il trouve 26 chaises libres. Par conséquent il a 26 possibilités de s'asseoir.*
- Der Student Nr. 2 kommt herein, Student Nr. 1 sitzt auf irgend einem Stuhl. Student Nr. 2 findet nur noch 25 freie Stühle vor. Somit hat er für sich nur noch 25 Sitzmöglichkeiten. Diese 25 Sitzmöglichkeiten hat er aber bei jeder Platzierung von Student Nr. 1, welcher sich auf 26 verschiedene Arten platzieren konnte. Zur ersten von Student Nr. 1 benutzten Möglichkeit hat Student Nr. 2 nun 25 Möglichkeiten, zur zweiten Möglichkeit von Student Nr. 1 hat Nr. 2 nun 25 Möglichkeiten, etc., zur letzten Möglichkeit von Student Nr. 1 hat Nr. 2 wiederum 25 Möglichkeiten. Zusammen haben beide also $26 \cdot 25$ Möglichkeiten. Die Anzahlen der Möglichkeiten multiplizieren sich!
- *L'étudiant no. 2 entre. L'étudiant no. 1 est assis sur une chaise quelconque. L'étudiant Nr. 2 trouve encore 25 chaises libres. Par conséquent il a seulement 25 possibilités de s'asseoir. Mais il a ces 25 possibilités de s'asseoir pour chaque position de l'étudiant no. 1, qui pouvait se placer sur 26 sièges différents. A la première possibilité utilisée par l'étudiant no. 1, l'étudiant no. 2 a maintenant 25 possibilités, pour la deuxième possibilité de l'étudiant no. 1, l'étudiant no. 2 a aussi 25 possibilités, etc, pour la dernière possibilité de l'étudiant no. 1, l'étudiant no. 2 a de nouveau 25 possibilités. En tout les deux ont ainsi $26 \cdot 25$ possibilités. Les nombres des possibilités se multiplient!*
- Der Student Nr. 3 kommt herein. Die Studenten Nr. 1 und Nr. 2 sitzen bereits. Student Nr. 3 findet nur noch 24 freie Stühle vor. Somit hat zu jeder der $26 \cdot 25$ Sitzmöglichkeiten der beiden ersten Studenten noch 24 Möglichkeiten. Zusammen haben sie also $26 \cdot 25 \cdot 24$ Möglichkeiten, da sich ja die Anzahlen der Möglichkeiten multiplizieren.
- *L'étudiant no. 3 entre. Les étudiants no. 1 et no. 2 sont déjà assis. L'étudiant Nr. 3 ne trouve que 24 chaises libres. Il a par conséquent pour chacune des $26 \cdot 25$ possibilités des premiers deux étudiants encore 24 possibilités. En tout les trois ont ainsi $26 \cdot 25 \cdot 24$ possibilités, parce que les nombres des possibilités se multiplient.*
- So geht es dann weiter. Schliesslich kommt der Student Nr. 25 herein. Er hat bei jeder der Sitzmöglichkeiten der vorher hineingegangenen Studenten noch 2 freie Plätze zur Auswahl. Total haben also die 25 ersten Studenten $26 \cdot 25 \cdot 24 \cdot \dots \cdot 3 \cdot 2$ Sitzmöglichkeiten.
- *Ainsi on avance. Finalement l'étudiant no. 25 entre. Pour chaque façon de se placer des étudiants qui sont déjà là il a encore 2 sièges de libres et par conséquent 2 possibilités de se placer. Totalement les 25 premiers étudiants ont ainsi $26 \cdot 25 \cdot 24 \cdot \dots \cdot 3 \cdot 2$ possibilités de se placer.*

- Endlich kommt der letzte Student mit der Nummer 26 herein. Er hat bei jeder der Sitzmöglichkeiten der andern Studenten noch einen freien Platz zur Auswahl. Total haben somit die 26 Studenten $26 \cdot 25 \cdot 24 \cdot \dots \cdot 3 \cdot 2 \cdot 1 = 26!$ Sitzmöglichkeiten.
- *Enfin le dernier étudiant, numéro 26, entre. Pour chaque façon de se placer des autres étudiants il ne lui reste qu'une chaise de libre, il n'a donc qu'une seule possibilité de se placer. Totalement les 26 étudiants ont par conséquent $26 \cdot 25 \cdot 24 \cdot \dots \cdot 3 \cdot 2 \cdot 1 = 26!$ possibilités de se placer.*

Bemerkung: • **Remarque:** Falls in einer Menge von Individuen jedes Element (Individuum) einen unterscheidbaren Namen trägt, so kann man die Elemente auch den „Namen nach“, d.h. alphabetisch ordnen, so wie in einem Lexikon: A...kommt vor B...etc., Aa...vor Ab...etc.. In einem solchen Fall spricht man von einer *lexikographischen Anordnung*.

- *Si dans un ensemble chaque élément (individu) porte un nom distinctif, on peut ranger les éléments "d'après les noms", c.-à.-d. de façon alphabétique, comme dans un lexique: A...vient avant B...etc, Aa...avant Ab...etc.. Dans un tel cas on parle d'une disposition lexicographique.*

Zum Nachdenken: • A réfléchir:

Falls die Klasse in 10 Sekunden einen Platzwechsel schafft, so braucht sie also $10 \cdot 26!$ Sekunden für alle Platzwechsel. Das sind $\frac{10 \cdot 26!}{60 \cdot 60 \cdot 24 \cdot 365}$ Jahre = 1.278310²⁰ Jahre. Vergleich: Das Alter des Universums bei der Urknalltheorie wird gegenwärtig auf ca. 1 bis 2 mal 10¹⁰ Jahre geschätzt⁵. Um die Sitzordnungen alle ohne Pause zu realisieren, bräuchte es also etwa 10¹⁰ mal soviel Zeit, wie das Universum alt ist!

- *Si la classe est capable d'exécuter un changement de place en 10 secondes, elle nécessite $10 \cdot 26!$ secondes pour tous les changements de place. Ça nous fait $\frac{10 \cdot 26!}{60 \cdot 60 \cdot 24 \cdot 365}$ ans = 1.278310²⁰ ans. Comparaison: L'âge de l'univers d'après la théorie du big bang est actuellement estimée à env. 1 à 2 fois 10¹⁰ ans⁶. Pour réaliser toutes les répartitions des places sans pauses, il faudrait donc 10¹⁰ fois l'âge de l'univers!*

Verallgemeinerung des Problems: • Généralisation du problème:

Statt mit 26 Studenten kann man das Problem gleich allgemein mit n Studenten lösen. In der Argumentation ist dann 26 durch n , 25 durch $n - 1$ etc. zu ersetzen. Man erhält schliesslich so total $(n!)$ Möglichkeiten, n Studenten auf n Plätze zu setzen.

- *Au lieu de résoudre le problème avec 26 étudiants on peut le résoudre généralement avec n étudiants. Dans l'argumentation il faut alors remplacer 26 par n , 25 par $n - 1$ etc.. Finalement on obtient totalement $(n!)$ possibilités de mettre n étudiants sur n places.*

Das abstrakte Problem — Le problème abstrait

Gegeben sei eine Menge $\mathcal{M}_n$ mit n Elementen, welche durchnummeriert sind mit den Nummern von 1 bis n . $\mathcal{M}_n$ entspricht der Menge der Studenten im vorherigen Beispiel. Dadurch hat man eine bijektive Zuordnung der nummerierten Elemente n_k zur Teilmenge der natürlichen Zahlen $\mathbf{N}_n = \{1, 2, 3, \dots, n\}$. (Damit hat man eine bijektive Funktion). Da die Zuordnung eineindeutig ist, können wir die n_k jeweils gerade durch k ersetzen, ohne das Problem zu verändern: $\mathcal{M} = \mathbf{N}_n = \{1, 2, 3, \dots, n\}$. Gesucht ist nun die Anzahl der Möglichkeiten, die Menge $\mathbf{N}_n = \{1, 2, 3, \dots, n\}$ auf sich selbst abzubilden, d.h. im obigen Problem die Menge der Nummern der Studenten $\mathbf{N}_n$ der Menge der Nummern der Stühle $\mathbf{N}_n$ zuzuordnen.

- *Soit donné un ensemble $\mathcal{M}_n$ avec n éléments, qui sont énumérés par les numéros de 1 jusqu'à n . $\mathcal{M}_n$ correspond à un ensemble d'étudiants dans l'exemple antérieur. On a ainsi un rapport bijectif des éléments numérotés n_k à un sous-ensemble $\mathbf{N}_n = \{1, 2, 3, \dots, n\}$ des nombres naturels. (On a ainsi une fonction bijective.) Comme le rapport est biunivoque, nous pouvons remplacer les n_k chaque fois par k , sans transformer le problème: $\mathcal{M} = \mathbf{N}_n = \{1, 2, 3, \dots, n\}$. Maintenant on cherche le nombre des possibilités d'appliquer l'ensemble $\mathbf{N}_n = \{1, 2, 3, \dots, n\}$ sur lui-même, c.-à.-d. dans le problème susdit appliquer l'ensemble des numéros des étudiants $\mathbf{N}_n$ à l'ensemble des numéros des chaises $\mathbf{N}_n$.*

⁵Die Fachleute streiten sich allerdings über diesen Wert. Je nach Wissensstand wird er laufend berichtigt.

⁶Les experts se disputent en effet au sujet de cette valeur. Elle est corrigée couramment d'après l'état des connaissances.

Sei $\sigma(k)$ bei einer solchen Zuordnung (im obigen Problem eine Sitzmöglichkeit) das Bild (oben die Stuhlnummer) von k (k entspricht oben der Nummer des Studenten). Dann wird also durch eine solche Zuordnung σ die Menge $\{1, 2, 3, \dots, n\}$ (oben die Menge der Studenten) der Menge $\{\sigma(1), \sigma(2), \sigma(3), \dots, \sigma(n)\}$ (oben die Stühle) zugeordnet. Schreibt man die Bilder $\sigma(k)$ unter die Urbilder k , so erscheint die durch σ ausgesonderte Relationsmenge in folgender Gestalt:

• Soit $\sigma(k)$ l'image (en haut le numéro des chaises) à une telle application (dans le problème susdit à une possibilité de s'asseoir) de k (en haut k correspond au numéro de l'étudiant). Alors par une telle application, on applique σ (en haut l'ensemble des étudiants) à l'ensemble $\{1, 2, 3, \dots, n\}$ (en haut les chaises). Si on écrit les images $\sigma(k)$ sous les originaux k , ainsi l'ensemble de relation défini par σ apparaît dans la forme suivante:

$$\begin{pmatrix} 1 & 2 & 3 & \dots & n \\ \sigma(1) & \sigma(2) & \sigma(3) & \dots & \sigma(n) \end{pmatrix}$$

Damit ist also eine Teilmenge von $\mathbf{N}_n \times \mathbf{N}_n$ gegeben, für die die Relation „Funktion σ “ zutrifft.

• Ainsi un sous-ensemble de $\mathbf{N}_n \times \mathbf{N}_n$ est donné pour lequel le rapport de "fonction σ " est valable.

Durch das folgende Schema wird daher eine neue Anordnung $\sigma(1), \sigma(2), \sigma(3), \dots, \sigma(n)$ der Elemente $1, 2, 3, \dots, n$ definiert.

• Par conséquent, par le schéma suivant, une nouvelle disposition $\sigma(1), \sigma(2), \sigma(3), \dots, \sigma(n)$ des éléments $1, 2, 3, \dots, n$ est définie.

$$\begin{pmatrix} 1 & 2 & 3 & \dots & n \\ \sigma(1) & \sigma(2) & \sigma(3) & \dots & \sigma(n) \end{pmatrix}$$

Wir sagen: • Nous disons:

Definition • Définition 3.3 (Permutation: • Permutation:)

Die Anordnung $\sigma(1), \sigma(2), \sigma(3), \dots, \sigma(n)$ der Elemente aus $\mathbf{N}_n$ heisst **Permutation $\mathcal{P}$** der Anordnung $(1, 2, 3, \dots, n)$ dieser Elemente.

• La disposition $\sigma(1), \sigma(2), \sigma(3), \dots, \sigma(n)$ des éléments de $\mathbf{N}_n$ s'appelle **permutation $\mathcal{P}$** de la disposition $(1, 2, 3, \dots, n)$ de ces éléments.

Um eine Permutation zu geben, können wir auch schreiben:

• Pour donner une permutation, nous pouvons aussi écrire:

$$P = \begin{pmatrix} 1 & 2 & 3 & \dots & n \\ \sigma(1) & \sigma(2) & \sigma(3) & \dots & \sigma(n) \end{pmatrix}$$

Die Reihenfolge der Spalten kann beliebig sein.

• Les collonnes se présentent dans un ordre quelconque.

Beispiel: • Exemple Durch die folgende Anordnung ist eine solche Permutation gegeben: • Par la disposition suivante, une telle permutation est donnée:

$$P = \begin{pmatrix} 1 & 2 & 3 & 4 & 5 \\ 4 & 1 & 5 & 3 & 2 \end{pmatrix}$$

1 wird auf 4, 2 auf 1 u.s.w.. abgebildet. Nun können wir unser Problem mit den Studenten und den Stühlen abstrakt und allgemein stellen:

• 1 est appliqué sur 4, 2 sur 1 etc.. Maintenant nous pouvons poser notre problème avec les étudiants et les chaises de façon abstraite et générale:

Problem • Problème 3.2

Permutationen ohne Wiederholung: • Permutations sans répétitions:

- Frage:** *Wieviele Permutationen $\mathcal{P}$ der Nummern $1, 2, \dots, n$ gibt es?*
- **Question:**
 - *Combien de permutations $\mathcal{P}$ des numéros $1, 2, \dots, n$ existent-ils?*
 - Oder anders gefragt: Wieviele Anordnungsmöglichkeiten der Zahlen $1, 2, \dots, n$ in einer Reihe gibt es?*
 - *Autrement: Combien de possibilités de dispositions des nombres $1, 2, \dots, n$ dans un rang existent-elles?*
 - Oder nochmals anders gefragt: Wieviele bijektive Funktionen $\mathbf{N}_n \mapsto \mathbf{N}_n$ gibt es?*
 - *Autrement encore: Combien de fonctions bijectives $\mathbf{N}_n \mapsto \mathbf{N}_n$ existent-elles?*

Symbole • **Symboles 1** : $P(n)$

Sei $P(n)$ = Anzahl Permutationen der Elemente von M_n der natürlichen Zahlen von 1 bis n .

- *Soit $P(n)$ = nombre des permutations des éléments de M_n des nombres naturels de 1 jusqu'à n .*

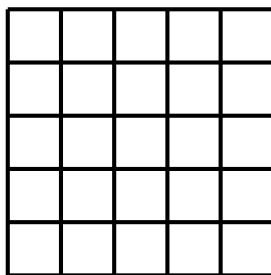
Nun wissen wir: • *Nous savons:*

Satz • **Théorème 3.2**

Permutationen ohne Wiederholung: • **Les permutations sans répétition:**

$$P(n) = n!$$

Abbildung 3.1: Teilflächen, verschieden zu färben ... • *Surfaces partielles, à colorer de manière différente...*



Beispiel: • *Exemple:* Wieviele Möglichkeiten gibt es, die in 3.1 gezeigten Teilflächen mit verschiedenen Farben zu färben? – Bei einer Färbung werden den 25 verschiedenen Flächen 25 verschiedene Farben zugeordnet. Statt Flächen und Farben kann man auch nur die Nummern 1 bis 25 betrachten. Man hat also eine bijektive Abbildung einer Menge $\mathcal{M}_{25}$ oder von $\mathbf{N}_{25}$ auf sich. Es wird also nach $P(25)$ = Anzahl Permutationen von $1, 2, 3, \dots, 25$ gefragt. Das gibt $25! \approx 1.55112 \cdot 10^{25}$. Wie lange hätte wohl einer, um alle Möglichkeiten auszuprobieren?

• *Combien de possibilités est-ce qu'il y a de colorer les différentes surfaces montrées dans 3.1 avec des couleurs différentes? – À une coloration, 25 couleurs différentes sont adjointes aux 25 surfaces différentes. Au lieu de surfaces et couleurs, on peut aussi considérer seulement les numéros de 1 jusqu'à 25. On a ainsi une application bijective d'un ensemble $\mathcal{M}_{25}$ ou de $\mathcal{M}_{25}$ sur soi-même. On cherche donc $P(25)$ = nombre de permutations de $1, 2, 3, \dots, 25$. Ça donne $25! \approx 1.55112 \cdot 10^{25}$. Combien de temps est-ce qu'il faudrait probablement pour exécuter toutes les possibilités?*

3.2.2 Permutationen mit Wiederholung — Permutations avec répétition

Paradigma — **Paradigme**

Problem • **Problème 3.3**

Vertauschungsmöglichkeiten von Briefen: • **Possibilités d'échange de lettres:**

- Situation:** *Ein Personalchef hat 20 verschiedene Briefe geschrieben. Davon sind 7 identische Kopien eines Informationsschreibens an eine Gruppe von Mitarbeitern, die andern 13 Briefe sind vertrauliche und persönliche Antworten in Lohnfragen anderer Mitarbeiter. Die zugehörigen 20 Couverts liegen ebenfalls bereit.*
- **Situation:** *• Un chef de personnel a écrit 20 lettres différentes. Il y a 7 copies identiques d'une lettre d'information pour un groupe de collaborateurs et 13 lettres concernant des réponses adressées d'autres collaborateurs concernant des questions de salaire. Les 20 enveloppes sont aussi prêtes.*
- Frage:** • **Question:** *Wieviele Möglichkeiten hat die Sekretärin, die Briefe zu verschicken, sodass für sie Probleme entstehen könnten?*
- *Combien de possibilités d'envoyer les lettres est-ce que la secrétaire a, de façon que pour elle des problèmes pourraient se poser?*

Lösung: • **Solution:** Wenn alle Briefe verschieden wären, so hätte sie $20!$ Möglichkeiten, die Briefe in Couverts zu stecken. Da nur eine Möglichkeit akzeptiert werden kann, führen dann $(20! - 1)$ Möglichkeiten zu Problemen.

Wenn nun 7 Briefe gleich sind, können diese 7 Briefe unter sich vertauscht werden, ohne dass ein Problem entsteht. Man kann das auf $7!$ Arten tun. Wenn nun X die Anzahl der Platzierungsmöglichkeiten der 13 verschiedenen Briefe in den 20 Couverts ist, so können zu jeder der X Möglichkeiten der verschiedenen Briefe die restlichen, gleichen Briefe auf $Y = 7!$ Arten unter sich vertauscht werden, ohne dass etwas passiert. Da das bei jeder der X Möglichkeiten der Fall ist, *multiplizieren sich die Anzahlen der Möglichkeiten zur Gesamtzahl der Möglichkeiten.* Andere Vertauschungsmöglichkeiten als die hier vorkommenden hat man nicht. Somit gilt: $20! = X \cdot Y = X \cdot 7!$ und damit $X = \frac{20!}{7!}$.

Die Anzahl unerwünschter Möglichkeiten ist somit $X - 1 = \frac{20!}{7!} - 1 = 482718652416000 - 1 \approx 4.82719 \cdot 10^{14}$.

• *Si toutes les lettres étaient différentes, elle auraient 20! possibilités de mettre les lettres dans les enveloppes. Comme une seule possibilité peut être acceptée, les $(20! - 1)$ autres possibilités causent des problèmes.*

Si 7 lettres sont maintenant les mêmes, ces 7 lettres peuvent être échangées entre elles sans qu'il y ait de problèmes. Cela on peut le faire de 7! manières différentes. Si maintenant X est le nombre de possibilités de placer les 13 lettres différentes dans le 20 enveloppes, pour chacune des X possibilités des lettres différentes les lettres identiques qui restent peuvent être échangées entre elles de $Y = 7!$ manières différentes sans qu'il se passe quelque chose d'embêtant. Comme c'est le cas pour chacune des X possibilités, les nombres des possibilités se multiplient au nombre total des possibilités. On n'a pas d'autres possibilités d'échange que celle qui sont mentionnées ici. Par conséquent il vaut: $20! = X \cdot Y = X \cdot 7!$ et par conséquent $X = \frac{20!}{7!}$.

Le nombre de possibilités indésirables est par conséquent $X - 1 = \frac{20!}{7!} - 1 = 482718652416000 - 1 \approx 4.82719 \cdot 10^{14}$.

Verallgemeinerung des Problems: • Généralisation du problème:

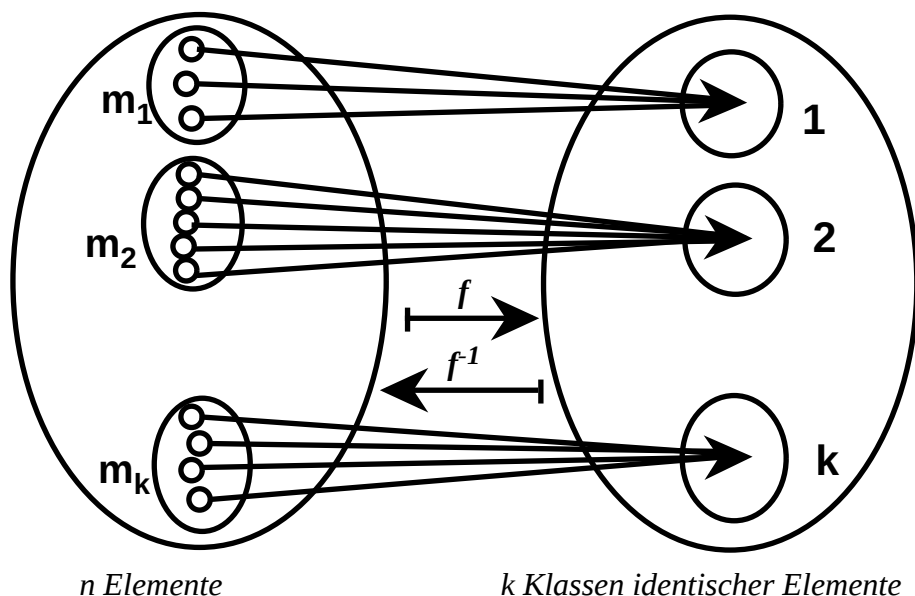
Wir gehen wieder von 20 Briefen aus, 7 davon gleich, die wir zur *Klasse 1* zusammenfassen. Weiter sei jetzt nochmals ein spezieller Brief da, zu welchem sich noch zwei gleiche finden. Diese 3 seien in einer *Klasse 2* zusammengefasst. Dann finden wir nochmals 4 gleiche, die in einer *Klasse 3* zusammengefasst werden. Sei nun Y_i die Anzahl der Vertauschungsmöglichkeiten der Briefe in der *Klasse i* unter sich und X wie vorhin die Anzahl der Vertauschungsmöglichkeiten der restlichen ungleichen Briefe. Dann gilt aus demselben Grunde wie oben:

• *Nous partons encore de 20 lettres dont 7 sont les mêmes qui sont groupées dans une classe 1. En plus on a une lettre spéciale, pour laquelle on trouve deux autres identiques. Ces 3 soient réunies dans une classe 2. Nous trouvons encore 4 identiques qui sont réunies dans une classe 3. Soit maintenant Y_i le nombre des possibilités d'échange des lettres entre elles dans la classe i et X comme en haut le nombre des possibilités d'échange des lettres inégales et restantes. Alors il vaut par la même raison comme en haut:*

$$20! = X \cdot Y_1! \cdot Y_2! \cdot Y_3! = X \cdot 7! \cdot 3! \cdot 4!, \text{ also } \bullet \text{ donc}$$

$$X = \frac{20!}{7! \cdot 3! \cdot 4!}$$

Abbildung 3.2: Anzahl möglicher Umkehrabbildungen f^{-1} ? • Possibilités d'applications inverses f^{-1} ?



• Esquisse: n éléments $\longmapsto$ n classes d'éléments identiques

Das führt uns auf folgendes allgemeinere Problem:

• Ça nous mène au problème général:

Gegeben: • **Donné:**

Total n Briefe, n Couverts, davon k Klassen je unter sich gleicher Briefe wie folgt:

• *Totalement n lettres, n enveloppes dont on a k classes de lettres identiques entre elles:*

Klasse₁: m_1 gleiche Briefe vom Typ 1 • Classe₁: m_1 lettres identiques du type 1,

Klasse₂: m_2 gleiche Briefe vom Typ 2 • Classe₂: m_2 lettres identiques du type 2,

⋮ ⋮

Klasse_k: m_k gleiche Briefe vom Typ k • Classe_k: m_k lettres identiques du type k

Gesucht: • **Trouver:** Anzahl Möglichkeiten $P_n(m_1, m_2, \dots, m_k)$, die Briefe in die Couverts zu platzieren.

• *Nombre de possibilités $P_n(m_1, m_2, \dots, m_k)$, de mettre les lettres dans les enveloppes.*

Symbole • **Symboles 2** : $P_n(m_1, m_2, \dots, m_k)$

$P_n(m_1, m_2, \dots, m_k)$ = Anzahl Möglichkeiten, die eben beschriebenen n Objekte (hier Briefe, wobei k Klassen mit je n_j gleichen Objekten darunter sind) auf n Plätze (hier Couverts) zu platzieren.

• $P_n(m_1, m_2, \dots, m_k)$ = nombre de possibilités, de placer les n objets qu'on vient de décrire (ici des lettres parmi lesquelles on trouve k classes avec n_j objets identiques entre eux) sur n places (ici des enveloppes).

Definition • **Définition 3.4 (Permutat. m. Wiederholung: • Permut. avec répétitions:)**

Gegeben sei eine Menge $\mathcal{M}_n$ mit n Elementen. Darin seien k Klassen mit je n_i gleichen Elementen (pro Klasse i) enthalten. Bei der Nummerierung der Elemente erhalten alle Elemente einer Klasse dieselbe

Nummer. Eine Permutation der Elemente von $\mathcal{M}_n$ nennen wir **Permutation mit Wiederholung**.

• Soit donné un ensemble $\mathcal{M}_n$ avec n éléments. Dans cet ensemble on trouve k classes avec n_i éléments identiques (par classe i). A une énumération des éléments, tous les éléments d'une classe reçoivent le même numéro. Nous appelons une permutation des éléments de $\mathcal{M}_n$ une **permutation avec répétitions**.

Wir wissen jetzt: • *Maintenant nous savons:*

Satz • Théorème 3.3 (Permutationen mit Wiederholung • Permut. avec répétition:) :

Anzahl der Permutationen mit Wiederholung: • Nombre de permutations avec répétitionss:

$$P_n(m_1, m_2, \dots, m_k) = \frac{n!}{m_1! \cdot m_2! \cdot m_k!}$$

Das abstrakte Problem — Le problème abstrait

Gegeben: • Donné: Eine Menge mit n Elementen, z.B. $\mathbf{R}_n = \{1, 2, 3, \dots, n\}$ sowie eine Menge mit k Elementen, z.B. $\mathbf{R}_k = \{1, 2, 3, \dots, k\}$, $n \geq k$. Man betrachte dann die möglichen Funktionen $f : \mathbf{R}_n \mapsto \mathbf{R}_k$ (ein Beispiel ist in Abb. 3.2 dargestellt).

• Un ensemble avec n éléments, par exemple $\mathbf{R}_n = \{1, 2, 3, \dots, n\}$ ainsi qu'un ensemble avec k éléments, z.B. $\mathbf{R}_k = \{1, 2, 3, \dots, k\}$. Puis on considère les fonctions possibles $f : \mathbf{R}_n \mapsto \mathbf{R}_k$ (un exemple est représenté dans image 3.2).

Gesucht: • Trouver: Anzahl möglicher Umkehrabbildungen $f^{-1} : \mathbf{R}_k \mapsto \mathbf{R}_n$. Dabei wird das erste Element (1 rechts im Bild) m_1 mal abgebildet, das zweite Element (2 rechts im Bild) m_2 mal u.s.w..

• Nombre d'applications inverses possibles $f^{-1} : \mathbf{R}_k \mapsto \mathbf{R}_n$. Pour cela le premier élément (1 à droite dans l'image) est appliqué m_1 fois, le deuxième élément (2 à droite dans l'image) m_2 fois etc..

Es werden also die k Klassen gleicher Elemente (gleiche Briefe im Paradigma) auf die n verschiedenen Elemente (Couverts im Paradigma) abgebildet. Die gesuchte Anzahl ist dann $P_n(m_1, m_2, \dots, m_k)$.

• Les k classes d'éléments identiques (lettres identiques dans le paradigme) sont ainsi appliquées sur les n éléments différents (enveloppes dans le paradigme). Le nombre qu'on cherche est donc $P_n(m_1, m_2, \dots, m_k)$.

Beispiel: • Exemple: Auf wieviele Arten lassen sich in einer Klasse mit 26 Studenten 5 Arbeitsgruppen mit 4, 5, 5, 6 und 6 Studenten bilden? Die Lösung ist:

• De combien de manières différentes est-ce qu'on peut former dans une classe de 26 étudiants 5 groupes de travail avec 4, 5, 5, 6 et 6 étudiants? La solution est:

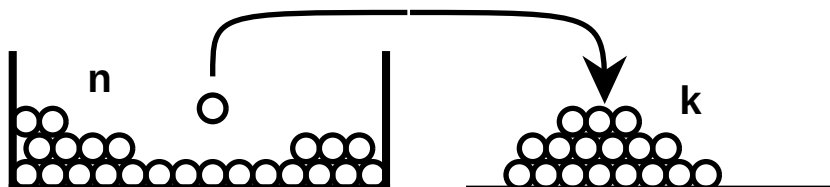
$$P_{26}(4, 5, 5, 6, 6) = \frac{26!}{4! \cdot 5!^2 \cdot 6!^2} = 2251024905729600 \approx 2.25102 \cdot 10^{15}$$

3.3 Auswahlprobleme mit und ohne Anordnung — Problèmes de sélection avec et sans rangement

3.3.1 Die Fragestellungen — Les questions

Kombinationen — Combinaisons

Problem • Problème 3.4 (Auswahlproblem: • Problème de sélection:)

Abbildung 3.3: Auswahlproblem, Kombinationen • *Problème de sélection, combinaisons*

Gegeben: • **Donné:** Eine Kiste mit n wohlunterscheidbaren Objekten, z.B. verschiedenfarbigen Kugeln. Aus der Kiste werden dann auf eine beliebige Art und Weise k Objekte herausgegriffen und nebeneinander aufgehäuft. (Vgl. Abb. 3.3.)

- Une caisse qui contient n objets bien distincts, par exemple des boules de différents couleurs. Dans la caisse, on choisit k objets de manière quelconque et on les accumule à côté. (Voir fig. 3.3.)

Frage: • **Question:** Auf wieviele Arten sind solche Haufenbildungen nebeneinander möglich?

- De combien de manières différentes peut-on former le tas à côté?

Wohlverstanden spielt bei der Haufenbildung die Anordnung der Objekte resp. der Kugeln keine Rolle. Dieses Problem lässt sich ohne viel Denkaufwand gleich abstrakt stellen. Die Kugeln in der Kiste bilden eine Menge $\mathcal{M}_n$, z.B. $\mathcal{M}_n = \mathbf{N}_n = \{1, 2, 3, \dots, n\}$. Herausgegriffen wird eine Teilmenge $\mathcal{M}_k \subseteq \mathcal{M}_n$, z.B. $\mathbf{N}_k = \{1, 2, 3, \dots, k\}$, $k \leq n$. Diese Teilmenge bildet den Haufen nebeneinander.

- La disposition des objets ne joue bien entendu aucun rôle à la disposition des objets resp. des boules. Il est possible de poser ce problème tout de suite abstraitement sans beaucoup de dépense de travail de cerveau. Les boules dans la caisse forment un ensemble $\mathcal{M}_n$, par exemple $\mathcal{M}_n = \mathbf{N}_n = \{1, 2, 3, \dots, n\}$. On choisit un sous-ensemble $\mathcal{M}_k \subseteq \mathcal{M}_n$, z.B. $\mathbf{N}_k = \{1, 2, 3, \dots, k\}$, $k \leq n$. Ce sous-ensemble forme le tas d'à côté.

Definition • Définition 3.5

Kombination ohne Wiederholung • Combinaison sans répétition

Eine solche Auswahl von k Elementen aus $\mathcal{M}_n$ heisst **Kombination k -ter Ordnung ohne Wiederholung** bei n Elementen, kurz: **Kombination k -ter Ordnung**.

- Un tel choix de k éléments dans $\mathcal{M}_n$ s'appelle **combinaison d'ordre k sans répétition** pour n éléments, brièvement: **Combinaison d'ordre k** .

Symbole • Symboles 3 (Anzahl Kombinationen: • Nombre de combinaisons:)

$C(k, n)$ = Anzahl Kombinationen k -ter Ordnung bei n Elementen.

- $C(k, n)$ = nombre de combinaisons d'ordre k pour n éléments.

Abstraktes Problem (Kombinationen ohne Wiederholung):

• Problème abstrait (combinaisons sans répétition):

Gegeben: • **Donné:** Eine Menge $\mathcal{M}_n$ mit n Elementen.

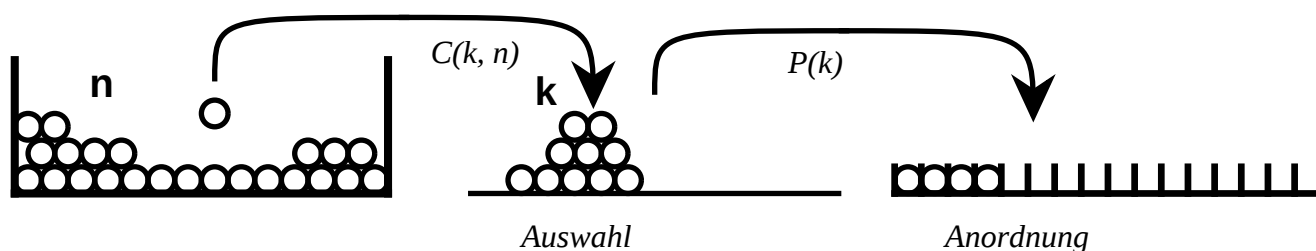
- Un ensemble $\mathcal{M}_n$ à n éléments.

Frage: • **Question:** $C(k, n) = ?$ D.h. wieviele Teilmengen mit genau k Elementen kann man bilden?

- $C(k, n) = ?$ C.-à.-d. combien de sous-ensembles peut-on former avec exactement k éléments?

Variationen — Arrangements

In Abb. 3.4 wird die Auswahl (Kombination) anschliessend noch angeordnet. Zwei solche Kombinationen mit denselben Elementen, aber verschiedener Anordnungen sind jetzt unterscheidbar. Man definiert da-

Abbildung 3.4: Relationsmenge, Abbildung • *Ensemble de relations et d'applications (choix, disposition)*

her:

• Dans la fig. 3.4 le choix (combinaison) est encore arrangé. On peut distinguer deux combinaisons de ce genre avec les mêmes éléments, mais de dispositions différentes. On définit par conséquent:

Definition • Définition 3.6

Variation ohne Wiederholung: • Arrangement sans répétition:

Werden die aus $\mathcal{M}_n$ ausgewählten Elemente (die Kombination also) noch angeordnet, so spricht man von einer Variation k -ter Ordnung ohne Wiederholung bei n Elementen. Kurz: Variation k -ter Ordnung.

• Si les éléments choisis dans $\mathcal{M}_n$ sont encore arrangés, on parle d'un arrangement d'ordre k sans répétition pour n éléments. Brièvement: Arrangement d'ordre k .

Symbole • Symboles 4 (Anzahl Variationen: • Nombre d'arrangements:)

$V(k, n) = \text{Anzahl Variationen } k\text{-ter Ordnung bei } n \text{ Elementen.}$

• $V(k, n) = \text{nombre d'arrangements d'ordre } k \text{ pour } n \text{ éléments.}$

Beispiel: • Exemple

Gegeben seien die Elemente a, b und c . Gesucht sind alle Kombinationen und Variationen 2-ter Ordnung.

• Soient donnés les éléments a, b et c . Trouver toutes les combinaisons et toutes les arrangements d'ordre 2.

Lösung: • Solution:

Kombinationen • Combinaisons:	$a b$	$a c$	$b c$:	3 Stück. • 3 pièces
Variationen: • Arrangements:	$a b$	$a c$	$b c$	
	$b a$	$c a$	$c b$:	6 Stück. • 6 pièces

Wiederholungen — Répétitions

Ersetzt man in der Vorratsmenge $\mathcal{M}_n$ jedes der Elemente e_i durch eine Menge E_i mit gleichen Elementen, die sich nur durch einen internen Index unterscheiden (z.B. $E_i = \{e_{i1}, e_{i2}, e_{i3}, \dots\}$), so wird es möglich, ein Element e_i mehrmals auszuwählen, wobei der interne Index nach der Auswahl wieder weggelassen werden kann⁷. Denselben Effekt erzielen wir, wenn wir nach der Auswahl eines Elementes eine identische Kopie dieses Elementes wieder zurücklegen. Wir stellen uns also vor, dass sich ein Element e_i bei seiner Auswahl dupliziert, sodass trotz Auswahl und Entfernung des Elements die Menge $\mathcal{M}_n$ unverändert bleibt. Ein Element wird also bei der Auswahl und Entfernung aus $\mathcal{M}_n$ sofort wieder in $\mathcal{M}_n$ nachgeliefert, etwa so wie bei einem bestimmten Artikel im Regal eines gut geführten Selbstbedienungsladens, wo die Regale immer sofort wieder aufgefüllt werden. Falls dieses Auffüllen, Duplizieren, Kopieren oder Zurücklegen beliebig oft möglich ist, so sagen wir, die Elemente in $\mathcal{M}_n$ seien *wiederholt auswählbar*. Wir definieren

⁷Der interne Index wird nur zur Bildung der „Mengen gleicher Elemente E_i “ gebraucht, die notwendig sind, um eine wiederholte Auswahl desselben Elements möglich zu machen.

nun:

- Si on remplace dans l'ensemble de réserve M_n chacun des éléments e_i par un ensemble E_i avec les mêmes éléments, qui ne se distinguent que par un indice interne (par exemple $E_i = \{e_{i1}, e_{i2}, e_{i3}, \dots\}$), ainsi il devient possible de choisir un élément e_i plusieurs fois. L'indice interne peut être omis après le choix⁸. Nous obtenons le même effet si nous mettons en réserve une copie identique de cet élément après le choix de l'élément. Nous nous imaginons donc qu'un élément e_i se duplique lors de son choix, et, que malgré qu'on ait choisi et enlevé l'élément, l'ensemble M_n reste inchangé. Un élément est donc fourni tout de suite dans M_n lors qu'on l'a choisi et enlevé de M_n . On peut comparer cela à la situation dans un supermarché où les marchandises sont remplacées sur les étagères au fur et à mesure qu'elles sont vendues. S'il est possible aussi souvent qu'on veut de remplir, copier ou remettre les éléments, nous disons que les éléments dans M_n sont répétitivement sélectionnables. Nous définissons maintenant:

Definition • Définition 3.7

Kombination und Variation mit Wiederholung: • **Combinaison et arrangement avec répétition:**

Sind bei der Bildung einer Kombination oder einer Variation die Elemente aus M_n wiederholt auswählbar, so spricht man von einer Kombination oder einer Variation mit **Wiederholung**.

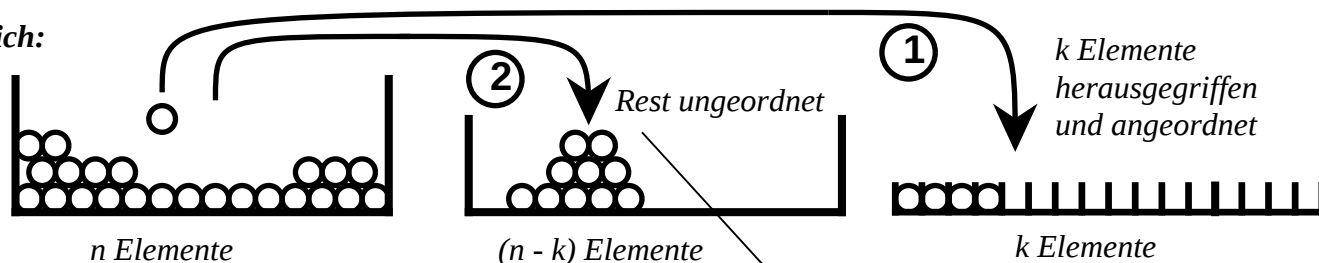
- Si à la formation d'une combinaison ou d'un arrangement les éléments de M_n sont sélectionnables de façon répétitive, nous parlons d'une combinaison ou d'un arrangement **avec répétition**.

Wir beginnen nun mit der Variation ohne Wiederholung:

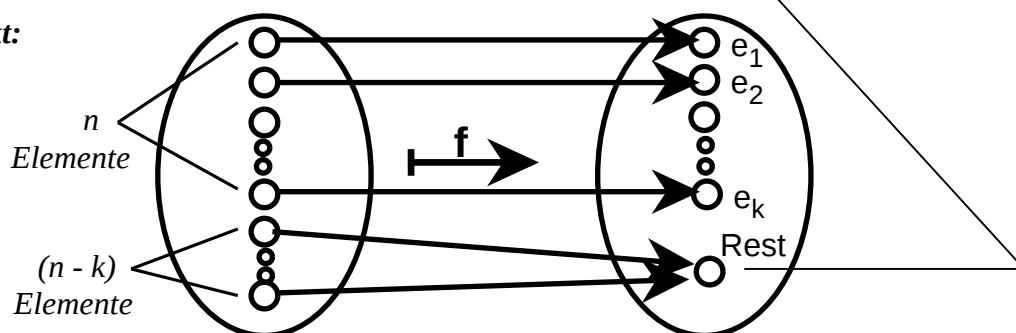
- Nous commençons maintenant avec l'arrangement sans répétition:

Abbildung 3.5: Variationen ohne Wiederholung • Arrangement sans répétition (image — abstrait, éléments, reste)

Bildlich:



Abstrakt:



⁸L'indice interne n'est utilisé que pour former l'ensemble d'éléments identiques E_i , qui sont nécessaires pour rendre possible un choix répété d'éléments identiques.

3.3.2 Variation ohne Wiederholung — Arrangement sans répétition

Aus n Elementen werden k Elemente herausgegriffen und angeordnet, ohne Wiederholung, so wie in Abb. 3.5 dargestellt. Dort wird z.B. das Element e_1 auf den Platz 1, e_2 auf den Platz 2 u.s.w.. gelegt. Alle $(n - k)$ nicht ausgewählten Elemente, der Rest also, kann man sich anschliessend in eine Kiste nebenan gelegt denken, auf einen Haufen also. Diese anschliessende Operation verändert die Anzahl Auswahl- und Anordnungsmöglichkeiten der ersten k Elemente nicht, denn diese Haufenbildung ist eine einzige, unabhängige Handlung, die nichts weiteres beiträgt. In dieser Restkiste nebenan spielt also die Anordnung der Elemente keine Rolle. Man unterscheidet diese Elemente demnach nicht, es ist egal, wie sie liegen. Daher bilden sie eine Klasse nicht unterschiedener, also gleicher Elemente, die auf nur eine einzige Art angeordnet werden können (da sie als nicht unterscheidbar gelten). Daher hat man folgendes Problem: Man hat n Elemente, k verschiedene und $n - k$ gleiche. Diese Elemente sind anzuordnen. Oder abstrakt: Man sucht die Anzahl der möglichen Umkehrfunktionen f^{-1} (vgl. Abb. 3.5). Das Problem haben wir aber bereits bei den Permutationen mit Wiederholung gelöst: Die Anzahl ist $P_n(n - k) = \frac{n!}{(n-k)!}$.

• *Dans n éléments on choisit k éléments qui sont disposés immédiatement, sans répétition d'éléments, à l'instar de fig. 3.5. Là par exemple l'élément e_1 est mis sur la place 1, e_2 sur la place 2 etc.. Ensuite on s'imagine que tous les $(n - k)$ éléments non-choisis, donc le reste, sont mis dans une caisse à part resp. sur un tas. Cette opération ne change pas le nombre de possibilités de choix, le nombre de possibilités de disposition des premiers k éléments, car cette formation de tas est une action unique et indépendante qui ne contribue rien à l'opération. Dans cette caisse de restes à part, la disposition des éléments ne joue pas de rôle. On ne distingue donc pas ces éléments, c'est égal comme ils sont disposés. Par conséquent ils forment une classe d'éléments non-distingués et donc une classe d'éléments égaux qui sont disposés d'une seule manière (parce qu'ils comptent comme non-distinctifs). Par conséquent on a le problème suivant: On a n éléments, k sont distinctifs et $n - k$ sont égaux. Ces éléments sont à arranger. Ou bien abstraitement: On cherche le nombre des fonctions inverses possibles f^{-1} (voir fig. 3.5). Ce problème a été résolu déjà à l'occasion des permutations avec répétition: Le nombre est $P_n(n - k) = \frac{n!}{(n-k)!}$.*

Satz • Théorème 3.4 (Variationen ohne Wiederholung: • Arrangements sans répétition:)

$$V(k, n) = P_n(n - k) = \frac{n!}{(n-k)!}$$

Beispiel: • Exemple:

Auf wieviele Arten kann man 20 verschiedene vorhandene Ferienjobs an 26 verschiedene Studenten verteilen, die alle einen solchen Job haben wollen, wenn diese Jobs nicht in Teiljobs aufteilbar sind?

Es handelt sich um die Auswahl 20 aus 26 mit anschliessender Zuordnung zu unterscheidbaren Studenten, d.h. Anordnung. Die Lösung ist somit:

• *De combien de manières est-ce qu'on peut distribuer 20 jobs de vacances différents et disponibles à 26 étudiants différents qui veulent avoir tous un semblable travail, si ces jobs ne sont pas divisibles dans des job partiels?*

Il s'agit du choix de 20 sur 26 avec un classement des étudiants non distinctifs, c.-à.-d. un arrangement. La solution est par conséquent:

$$V(20, 26) = \frac{26!}{(26 - 20)!} = \frac{26!}{(6)!} = 67215243521100939264000000 \approx 6.72152 \cdot 10^{25}$$

Ein Spezialfall: • Un cas spécial: $V(n, n) = P_n(n - n) = P_n(0) = P(n)$

$\leadsto$ Permutation ohne Wiederholung! • *Permutation sans répétition!*

3.3.3 Kombination ohne Wiederholung — Combinaison sans répétition

Die Formel — La formule

Auf Seite 43 haben wir gesehen, dass sich bei Aussonderung einer Teilmenge gleicher Elemente die Anzahl der Möglichkeiten multiplikativ verhalten. Da war $20! = X \cdot Y = X \cdot 7!$. Die gleiche Situation finden wir

beim Übergang von den Kombinationen zu den Variationen: Eine Variation (k Elemente aus n Elementen) entsteht aus einer Kombination durch Anordnung der k ausgewählten Elemente. Dazu hat man $P(k) = k!$ Möglichkeiten. Es gilt also:

• A la page 43 nous avons vu qu'à une sélection d'un sous-ensemble de mêmes éléments le nombre des possibilités se comporte de façon multiplicative. On y a trouvé: $20! = X \cdot Y = X \cdot 7!$. Nous trouvons la même situation au passage des combinaisons à l'arrangement: Un arrangement (k éléments de n éléments) peut être obtenu d'une combinaison par la disposition des k éléments choisis. On y a $P(k) = k!$ possibilités. Il vaut donc:

Lemma • Lemme 3.1 (Variationen und Kombination: • Arrangements et combinaison:)

$$V(k, n) = C(k, n) \cdot P(k), \text{ also } \frac{n!}{(n-k)!} = C(k, n) \cdot k!$$

Daraus folgt: • Il en suit:

Satz • Théorème 3.5 (Kombination ohne Wiederholung • Combinaison sans répétition:)

$$C(k, n) = \frac{n!}{k! \cdot (n-k)!} = \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}$$

Das Beispiel Zahlenlotto „6 aus 45“: • L'exemple du jeu de loto "6 de 45":

Auf wieviele Arten kann man 6 verschiedene Zahlen aus den 45 ersten natürlichen Zahlen auswählen? Hier handelt es sich um eine typische Frage nach der Anzahl Kombinationen $C(6, 45)$. Diese ist gleich:

• De combien de manières différentes est-ce qu'on peut choisir 6 nombres différents dans les 45 premiers nombres naturels? Ici, il s'agit d'une question typique concernant le nombre des combinaisons $C(6, 45)$. Celle-ci est égale à:

$$\frac{45!}{6! \cdot (45-6)!} = \frac{45!}{6! \cdot (39)!} = 8145060 \approx 8.14506 \cdot 10^6$$

Binomialkoeffizienten — Coefficients binomiaux

Multipliziert man das Binom $(a+b)^n = \overbrace{(a+b) \cdot (a+b) \cdot \dots \cdot (a+b)}^{n \text{ Faktoren}}$ nach den Regeln des Distributivgesetzes aus, so entstehen lauter Summanden der Form $m_k \cdot a^k \cdot b^{n-k}$ mit $0 \leq k \leq n$ und $m, k, n \in \mathbf{N}_0$. Beim Ausmultiplizieren nimmt man der Reihe nach aus jedem Faktor $(a+b)$ einen der Summanden a oder b und multipliziert diese Faktoren zu einem Produkt $a^k \cdot b^{n-k}$. Falls man in jedem Summanden a und nie b nimmt, entsteht $a^n \cdot b^0$. Falls man in j Summanden a und folglich in $n-j$ Summanden b nimmt, entsteht $a^j \cdot b^{(n-j)}$. Dabei gibt es hier verschiedene Möglichkeiten, das a oder das b auszuwählen: Man kann z.B. im ersten Faktor a , im zweiten b , im dritten wieder a wählen etc., man kann aber auch zuerst b , dann a und dann wieder a wählen etc.. m_k ist die Anzahl der Möglichkeiten, a in genau k Faktoren $(a+b)$ und b in genau $n-k$ Faktoren zu wählen. Es ist dann:

• Si on multiplie le binôme $(a+b)^n = \overbrace{(a+b) \cdot (a+b) \cdot \dots \cdot (a+b)}^{n \text{ facteurs}}$ d'après les règles de la loi distributive,

on n'obtient que des termes additionnels de la forme $m_k \cdot a^k \cdot b^{n-k}$ avec $0 \leq k \leq n$ et $m, k, n \in \mathbf{N}_0$. A la multiplication on prend selon le rang de chaque facteur $(a+b)$ un des termes additionnels a ou b et on les multiplie en obtenant un produit $a^k \cdot b^{n-k}$. Si on prend dans chaque terme additionnel a et jamais b , on obtient $a^n \cdot b^0$. Si on prend a dans j termes additionnels et b dans $n-j$ termes additionnels, on obtient $a^j \cdot b^{(n-j)}$. A cette occasion il existe plusieurs possibilités de choisir le a ou le b . Par exemple on peut choisir a dans le premier facteur, b dans le deuxième facteur, de nouveau a dans le troisième facteur etc., mais on peut aussi choisir d'abord b , après a et alors encore une fois a etc... m_k est le nombre des

possibilités de choisir a dans exactement k facteurs $(a + b)$ et b dans exactement $n - k$ facteurs. Il vaut donc :

$$(a + b)^n = \sum_{k=0}^n m_k \cdot a^k \cdot b^{n-k}$$

Wie gross ist nun m_k ? — Hier handelt es sich um ein Auswahlproblem: Auf wieviele Arten kann man aus den n Faktoren $(a+b)$ k Faktoren auswählen und dort jeweils den Anteil a (und folglich in den restlichen $n-k$ Faktoren jeweils den Anteil b) nehmen? Diese Frage ist äquivalent mit der einfacheren Frage: Auf wieviele verschiedene Arten kann man aus n Elementen (Faktoren $(a+b)$) jetzt k Elemente auswählen? Die Antwort ist nun einfach: $m_k = C(k, n)$. m_k hat einen Namen:

• Quelle est maintenant la grandeur de m_k ? – Ici, il s'agit d'un problème de choix: De combien de manières différentes est-ce qu'on peut choisir entre les n facteurs $(a+b)$ k facteurs et y prendre la partie a (et par conséquent prendre dans les $n-k$ facteurs restants chaque fois la partie b)? Cette question est équivalente à la question plus simple: De combien de manières différentes est-ce qu'on peut maintenant choisir entre n éléments (facteurs $(a+b)$) k éléments? La réponse est maintenant très simple: $m_k = C(k, n)$. m_k porte un nom:

Definition • Définition 3.8 (Binomialkoeffizient: • Coefficient binomial:) :

m_k heisst Binomialkoeffizient. • m_k s'appelle Coefficient binomial.

Symbole • Symboles 5 (Binomialkoeffizient: • *Coéfficient binomial:*) $m_k := \binom{n}{k}$

Wir wissen jetzt: • *Nous savons maintenant:*

Satz • Théorème 3.6 (Binomischer Lehrsatz • Théorème binomial:)

$$(a+b)^n = \sum_{k=0}^n m_k \cdot a^k \cdot b^{n-k} = \sum_{k=0}^n \binom{n}{k} \cdot a^k \cdot b^{n-k}$$

Binomialkoeffizienten kann man bekanntlich im *Pascalschen Dreieck* ablesen:

- On peut lire les coefficients binomiaux dans le triangle de Pascal:

Pascalsches Dreieck: • **Triangle de Pascal:**

[illegible]

Die vertikale Position ist n , die horizontale k . Die Numerierung beginnt jeweils mit 0. So liest man z.B. ab: $\binom{4}{1} = 4$ und $\binom{4}{2} = 6$.

• La position verticale est n , la position horizontale k . Le numérotage commence chaque fois par 0. Ainsi on peut lire par exemple: $c_{(1)}^{(4)} = 4$ und $c_{(2)}^{(4)} = 6$.

Für die Binomialkoeffizienten kann man mit Hilfe von $\binom{n}{k} = C(k, n) = \frac{n!}{k! \cdot (n-k)!} = \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}$ sowie mit dem Prinzip der vollständigen Induktion⁹ leicht die folgenden Gesetze beweisen:

- Pour les coefficients binomiaux, on peut prouver à l'aide de $\binom{n}{k} = C(k, n) = \frac{n!}{k! \cdot (n-k)!} = \frac{n \cdot (n-1) \cdot \dots \cdot (n-k+1)}{1 \cdot 2 \cdot \dots \cdot k}$ ainsi qu'avec le principe de l'induction complète¹⁰ facilement les lois suivantes:

Satz • Théorème 3.7

Einige Eigenschaften der Binomialkoeffizienten: • Quelques qualités des coefficients binomiaux:

⁹Vgl. Zahlenlehre

¹⁰Voir théorie des nombres

$$\begin{array}{ll}
1) & \binom{n}{k} = \binom{n}{n-k} \\
3) & 2^n = \sum_{k=0}^n \binom{n}{k} \\
5) & \sum_{s=0}^{n-1} \binom{k+s}{k} = \binom{n+k}{k+1} \\
2) & \binom{n-1}{k-1} + \binom{n-1}{k} = \binom{n}{k} \\
4) & \sum_{k=0}^r \binom{p}{k} \cdot \binom{q}{r-k} = \binom{p+q}{r} \\
6) & \sum_{k=0}^p \binom{p}{k}^2 = \binom{2p}{p}
\end{array}$$

Z.B. die Formel $2^n = \sum_{k=0}^n \binom{n}{k}$ ergibt sich aus $2^n = (1+1)^n = \sum_{k=0}^n \binom{n}{k} \cdot 1^k \cdot 1^{n-k} = \sum_{k=0}^n \binom{n}{k}$ mit Hilfe des binomischen Lehrsatzes.

• *Par exemple la formule $2^n = \sum_{k=0}^n \binom{n}{k}$ est obtenue par $2^n = (1+1)^n = \sum_{k=0}^n \binom{n}{k} \cdot 1^k \cdot 1^{n-k} = \sum_{k=0}^n \binom{n}{k}$ à l'aide du théorème binomial.*

3.3.4 Variation mit Wiederholung — Arrangement avec répétition

Die Formel — La formule

Die *Variation mit Wiederholung* ist auf Seite 48 erklärt worden. Die Formel für die Anzahl Variationen mit Wiederholung hingegen müssen wir noch erarbeiten. Dazu verwenden wir folgendes Symbol:

• *L'arrangement avec répétition a été expliqué à la page 48. La formule pour le nombre d'arrangements avec répétition par contre doit encore être élaborée. Pour cela nous utilisons le symbole suivant:*

Symbole • Symboles 6 : $\bar{V}(k, n)$

$\bar{V}(k, n)$ = Anzahl Variationen mit Wiederholung bei einer Auswahl von k Elementen aus einem Vorrat mit n verschiedenen Elementen, die alle wiederholbar sind.

• $\bar{V}(k, n)$ = nombre d'arrangements avec répétition pour un choix de k éléments dans une réserve avec n éléments différents, qui tous peuvent être répétés.

Herleitung der Formel: • Déduction de la formule:

Wir betrachten die k nummerierten Plätze, auf denen die auszuwählenden Elemente anzuordnen sind (vgl. Abb. 3.5 oben links im Bild). Da wir jedes der n Elemente im Vorrat auswählen können, hat man n Möglichkeiten, den 1. Platz zu besetzen. Bei der Auswahl für den 2. Platz hat man aber wieder n Elemente im Vorrat zur Auswahl, da wegen der Wiederholbarkeit wieder jedes Element vorhanden ist und gewählt werden kann: Zu jeder der n Möglichkeiten für den 1. Platz hat man n Möglichkeiten für den 2. Platz, total also jetzt $n \cdot n = n^2$ Möglichkeiten. Genauso geht es für den 3. Platz: Zu jeder der n^2 Möglichkeiten für die Plätze 1 und 2 hat man n Möglichkeiten für den 3. Platz, total also jetzt $n^2 \cdot n = n^3$ Möglichkeiten. So fährt man fort: Für die Besetzung der ersten 4 Plätze hat man n^4 Möglichkeiten, für die Besetzung der ersten 5 Plätze n^5 Möglichkeiten und schliesslich für die Besetzung der ersten k Plätze hat man n^k Möglichkeiten. Wir haben somit den Satz:

• *Nous considérons les k places numérotés sur lesquelles les éléments à choisir sont ordonnés (voir fig. 3.5 en haut à gauche dans l'image). Comme nous pouvons choisir chacun des n éléments dans la réserve, on a n possibilités d'occuper la 1ère place. Pour la 2ème place on a de nouveau n éléments dans la réserve à disposition pour le choix; à cause de la possibilité de répétition chaque élément existe toujours et peut être choisi: Pour chacune des n possibilités pour la 1ère place on a n possibilités pour la 2ème place, totalement donc $n \cdot n = n^2$ possibilités. Également pour la 3ème place: Pour chacun des n^2 possibilités pour les places 1 et 2 on a n possibilités pour la 3ème place, totalement donc $n^2 \cdot n = n^3$ possibilités. On continue ainsi: Pour l'occupation des premières 4 places on a n^4 possibilités, pour l'occupation des premières 5 places n^5 possibilités et finalement pour l'occupation des premières k places on a n^k possibilités. Par conséquent nous avons le théorème:*

Satz • Théorème 3.8 (Variationen mit Wiederholung: • Arrangement avec répétitions:)

$$\bar{V}(k, n) = n^k$$

Beispiel: • Exemple:

Auf wieviele Arten können 26 (unterscheidbare) Studenten sich in 12 verschiedene Kurse einschreiben, wenn jeder Kurs 26 Plätze offen hat, also keine Platzbeschränkung besteht?

- De combien de possibilités différentes est-ce que 26 étudiants (qu'on peut distinguer) peuvent s'inscrire dans 12 cours différents, si chaque cours offre 26 places, c.à.d. s'il n'y a pas de limites aux places?

Lösung: • **Solution:**

Der erste Student hat 12 Möglichkeiten, sich in einen Kurs einzuschreiben. Zu jeder dieser Möglichkeiten des ersten Studenten hat der zweite auch 12 Möglichkeiten, sich in einen Kurs einzuschreiben. Beide zusammen haben also 12^2 Möglichkeiten. Für den dritten, vierten etc. Studenten geht das auch so: Jeder hat die 12 Möglichkeiten, und die Möglichkeiten multiplizieren sich. Es handelt sich um eine Variation mit Wiederholung. Total gibt es $\bar{V}(k, n) = \bar{V}(26, 12) = 12^{26} = 11447545997288281555215581184 \approx 1.14475 \cdot 10^{28}$ Möglichkeiten.

• Le premier étudiant a 12 possibilités de s'inscrire dans un cours. Pour chacune de ces possibilités du premier étudiant le deuxième a aussi 12 possibilités de s'inscrire dans un cours. Les deux ensemble ont 12^2 possibilités. Pour le troisième, quatrième etc. étudiant ça fonctionne aussi d'après le même schéma: Chacun a les 12 possibilités, et les possibilités se multiplient. Il s'agit d'un arrangement avec répétitions. Totalemment il y a $\bar{V}(k, n) = \bar{V}(26, 12) = 12^{26} = 11447545997288281555215581184 \approx 1.14475 \cdot 10^{28}$ possibilités.

Merke: Aus diesem Beispiel ersieht man, dass $k > n$ sein kann.

- **A retenir:** Par cet exemple, on voit que k peut être plus grand que n : $k > n$.

Anwendung: Die Mächtigkeit der Potenzmenge — Application: Puissance de l'ensemble de parties

Die Potenzmenge ist bekanntlich die Menge aller Teilmengen.

- L'ensemble de parties est comme chacun sait l'ensemble de tous les sous-ensembles.

Problem • **Problème 3.5**

Mächtigkeit der Potenzmenge: • **La puissance de l'ensemble de parties**

Gegeben: • **Donné:** Eine Menge $\mathcal{M}$ mit n Elementen.

- Un ensemble $\mathcal{M}$ à n éléments.

Frage: • **Question:** Wieviele Teilmengen hat $\mathcal{M}$?

- Combien d'ensembles partiels $\mathcal{M}$ a-t-il?

Lösung: • **Solution:**

$\binom{n}{k} = C(k, n)$ ist bekanntlich die Anzahl Teilmengen mit k Elementen, denn hier handelt es sich ja um das typische Auswahlproblem. Nun kann man eine oder mehrere Teilmengen mit 0 (leere Menge), 1, 2, ..., n Elemente wählen. Total hat man also:

• $\binom{n}{k} = C(k, n)$ est comme chacun sait le sous-ensembles avec k éléments, car ici, il s'agit d'un problème de choix typique. Maintenant on peut choisir un ou plusieurs sous-ensembles avec 0 (quantité vide), 1, 2, ..., n éléments. Totalemment on a donc:

$$\binom{n}{0} + \binom{n}{1} + \binom{n}{2} + \dots + \binom{n}{n} = \sum_{k=0}^n \binom{n}{k} = \sum_{k=0}^n \binom{n}{k} \cdot 1^k \cdot 1^{n-k} = (1+1)^n = 2^n.$$

Satz • **Théorème 3.9 :**

Mächtigkeit der Potenzmenge: • **Puissance de l'ensemble de parties**

Die Potenzmenge einer Menge mit n Elementen hat 2^n Elemente.

- L'ensemble de parties d'un ensemble qui contient n éléments possède 2^n éléments.

Eine Menge mit n Elementen hat also genau 2^n Teilmengen.

- Donc un ensemble avec n éléments contient exactement 2^n sous-ensembles.

3.3.5 Kombination mit Wiederholung — Combinaison avec répétition

Hier sollen aus einer Menge mit n Elementen k Elemente ausgewählt werden, wobei jedes ausgewählte Element bei der Auswahl in der Menge dupliziert wird resp. nachgeliefert wird, so dass die Menge trotz Auswahl immer aus denselben Elementen besteht. Wie gross ist die Anzahl Auswahlmöglichkeiten?

• *Ici il faut choisir k éléments dans un ensemble avec n éléments. Chaque élément choisi se duplique dans l'ensemble de façon que l'ensemble reste toujours le même malgré le choix. Quel est le nombre des options quant au choix?*

Für die Berechnung dieser Anzahl ist es unwesentlich, ob die Menge $\mathcal{M}_n$ aus Kugeln, Losen oder Zahlen etc. besteht, d.h. welcher Natur die Elemente sind. Wir dürfen daher annehmen, es handle sich um die natürlichen Zahlen von 1 bis n : $\mathcal{M}_n = \{1, 2, 3, \dots, n\}$. Wenn wir jetzt k Elemente (d.h. Zahlen) auswählen, so wollen wir diese immer ihrer Grösse nach aufreihen, statt sie bloss „auf einen Haufen zu legen“. Wir reden hier von der *Standardanordnung*. Eine solche Auswahl $\{e_1, e_2, \dots, e_k\}$ wird also immer in der Anordnung $e_1 \leq e_2 \leq \dots \leq e_k$ präsentiert. Dadurch wird die Anzahl der Auswahlmöglichkeiten ja nicht verändert.

• *Pour le calcul de ce nombre, il n'est pas essentiel si l'ensemble $\mathcal{M}_n$ consiste en boules, en billets de lotterie ou en nombres, c.-à.-d. quelle est la nature des éléments. Nous pouvons supposer par conséquent qu'il s'agit de nombres naturels de 1 jusqu'à n : $\mathcal{M}_n = \{1, 2, 3, \dots, n\}$. Si nous choisissons maintenant k éléments (c.-à.-d. des nombres), nous voulons les ranger l'un à côté de l'autre au lieu de les "mettre sur un tas". Nous parlons ici du rangement standard (configuration). Nous présentons donc un tel choix $\{e_1, e_2, \dots, e_k\}$ toujours dans une disposition $e_1 \leq e_2 \leq \dots \leq e_k$. Par cela le nombre des options ne change pas.*

Wie ist nun dem Problem der Wiederholungen beizukommen? Die Idee, aus $k \cdot n$ Elementen auszuwählen, führt zu keinem Resultat, da die Elemente einer Auswahlmenge dann auf verschiedene Weise gewonnen werden können, was fälschlicherweise die Anzahl Auswahlmöglichkeiten erhöht. So geht es also nicht. Um der Sache beizukommen, muss man etwas weiter ausholen:

• *Comment résoudre le problème des répétitions? L'idée de choisir parmi $k \cdot n$ éléments ne mène à aucun résultat parce que les éléments d'un ensemble de choix peuvent être obtenus de façon différente ce qui augmente faussement le nombre des possibilités de choix. Donc ça ne va pas de cette manière. Pour venir à bout de la chose on doit aller chercher plus loin:*

Wir führen dazu $k - 1$ neue Elemente $J_1, J_2, \dots, J_{k-1}$ ein und fügen diese der Menge $\mathcal{M}_n$ an. So erhalten wir eine neue Menge $\mathcal{M}_n^{k-1} = \{1, 2, 3, \dots, n, J_1, J_2, \dots, J_{k-1}\}$ mit $n + k - 1$ Elementen. Die neu gültige Standardanordnung entspreche der hier gegebenen Aufzählung der Elemente: Die J_i werden hinten den Nummern nach angefügt. Dabei gelte für die Elemente J_i die folgende Interpretation: J_i ist eine Vorschrift oder Funktion, die auf jenen ausgewählten Standardanordnungen operiert, in denen sie selbst allenfalls vorkommt. Die durch J_i gegebene Vorschrift lautet: Ersetze das Symbol J_i in einer ausgewählten Standardanordnung durch das Element e_i derselben Auswahl, nachdem alle J_p mit $p < i$ schon ersetzt sind. Führt man alle diese Ersetzungen durch, so erhält man aus einer *primären Auswahl* die *Endstandardanordnung*. Da k Elemente auszuwählen sind, es aber nur $k - 1$ Elemente J_i gibt, kommt in einer Standardauswahl immer mindestens ein Element $e_j \in \mathcal{M}_n$ vor, in unserem Falle eine der gegebenen natürlichen Zahlen $1, 2, 3, \dots, n$. J_i bewirkt somit immer eine Ersetzung durch ein weiter vorne vorkommendes Element in der Standardanordnung, also eine Duplikation. Da so jedes Element einmal ausgewählt und dann noch durch die J_i maximal $k - 1$ mal dupliziert werden kann, besteht die Möglichkeit, dass jedes Element von $\mathcal{M}_n$ dann k mal in der Endstandardanordnung vorkommen kann. Auf diese Art können alle Kombinationen mit Wiederholung gewonnen werden.

• *A cette intention nous introduisons $k - 1$ nouveaux éléments $J_1, J_2, \dots, J_{k-1}$, les incluons dans l'ensemble $\mathcal{M}_n$. Ainsi nous recevons un nouveau ensemble $\mathcal{M}_n^{k-1} = \{1, 2, 3, \dots, n, J_1, J_2, \dots, J_{k-1}\}$ à $n + k - 1$ éléments. La disposition standard nouvellement valable correspond à l'énumération des éléments donnée ici: Les J_i sont ajoutés derrière d'après les numéros. Pour les éléments J_i l'interprétation suivante est valable: J_i est une prescription ou fonction qui opère sur les dispositions standard choisies,*

dans lesquelles on les trouve elles-mêmes. La prescription donnée par J_i dit: Remplacer le symbole J_i dans une disposition standard choisie par l'élément e_i du même choix après avoir remplacé tous les J_p par $p < i$. Si on effectue tous ces remplacements, on obtient d'un choix primaire la disposition finale standard. Comme il faut choisir k éléments et comme il n'existent que $k - 1$ éléments J_i , dans un choix standard on trouve toujours au moins un élément $e_j \in \mathcal{M}_n$, dans notre cas un des nombres naturels $1, 2, 3, \dots, n$. J_i effectue par conséquent toujours un remplacement par un élément qui est situé plus à l'avant dans la disposition standard, donc une duplicata. Comme chaque élément peut être choisi une fois ainsi et après peut être dupliqué par un J_i au maximum $k - 1$ fois, il existe la possibilité que chaque élément de $\mathcal{M}_n$ peut se trouver donc k fois dans la disposition standard finale. De cette façon peuvent être obtenues toutes les combinaisons avec répétition.

Beispiel: • Exemple:

Gegeben sei $\mathcal{M}_7 = \{1, 2, 3, \dots, 7\}$. Daraus sollen 5 Elemente mit Wiederholung ausgewählt werden. Es ist dann $\mathcal{M}_7^{5-1} = \mathcal{M}_7^4 = \{1, 2, 3, \dots, 7, J_1, J_2, J_3, J_4\}$.

• Soit donné $\mathcal{M}_7 = \{1, 2, 3, \dots, 7\}$. Dans cet ensemble il faut choisir 5 éléments avec répétitions. Il vaut donc: $\mathcal{M}_7^{5-1} = \mathcal{M}_7^4 = \{1, 2, 3, \dots, 7, J_1, J_2, J_3, J_4\}$.

Wählt man z.B. $(1, 5, 7, J_1, J_4)$ (in Standardanordnung), so wird wie folgt ersetzt: Zuerst $J_1 \mapsto 1$ (der Index 1 ist kleiner als der Index 4). Das ergibt $(1, 5, 7, 1, J_4)$ in Nicht-Standardanordnung und $(1, 1, 5, 7, J_4)$ in neuer Standardanordnung. Dann wird ersetzt $J_4 \mapsto 7$, was zur Standardanordnung $(1, 1, 5, 7, 7)$ führt.

• Si on choisit par exemple $(1, 5, 7, J_1, J_4)$ (dans la disposition standard), il faut remplacer comme suit: D'abord $J_1 \mapsto 1$ (l'indice 1 est plus petit que l'indice 4). Ça donne $(1, 5, 7, 1, J_4)$ dans la disposition non-standard et $(1, 1, 5, 7, J_4)$ dans la nouvelle disposition standard. Alors on remplace $J_4 \mapsto 7$ ce qui mène à la disposition standard $(1, 1, 5, 7, 7)$.

Ähnlich führt die Auswahl $(4, J_1, J_2, J_3, J_4)$ nach allen Ersetzungen zur Standardanordnung $(4, 4, 4, 4, 4)$.

• Semblablement le choix $(4, J_1, J_2, J_3, J_4)$ mène à la disposition standard $(4, 4, 4, 4, 4)$ après tous les remplacements.

Bei der Auswahl von 6 Elementen aus $\mathcal{M}_8$ führt die primäre Auswahl $(2, 3, 7, 8, J_2, J_4)$ auf die Endstandardanordnung $(2, 3, 3, 7, 7, 8)$.

• Au choix de 6 éléments dans $\mathcal{M}_8$ le choix primaire mène à la disposition standard finale $(2, 3, 7, 8, J_2, J_4)$.

Diese Beispiele machen klar, dass eine primäre Auswahl eindeutig einer Endstandardanordnung entspricht. Die Anzahl der auswählbaren primären Anordnungen ist gleich der Anzahl der Endstandardanordnungen, in welchen alle Elemente bis zu k mal wiederholt vorkommen können. Um $\bar{C}(k, n)$ zu finden, muss man also die Anzahl der primär auswählbaren Standardanordnungen bestimmen. Dort werden k Elemente aus den $n + k - 1$ Elementen $1, 2, 3, \dots, n, J_1, J_2, \dots, J_{k-1}$ ausgewählt. Daher ist $\bar{C}(k, n) = C(k, n + k - 1)$. Somit hat man:

• Ces exemples montrent qu'un choix primaire correspond clairement à une disposition standard finale. Le nombre des dispositions primaires et sélectionnables est égal au nombre des dispositions standard finales, dans lesquelles tous les éléments figurent répétés jusqu'à k fois. Pour trouver $\bar{C}(k, n)$, on doit donc trouver le nombre des dispositions standard primaires sélectionnables. Là, k éléments sont choisis entre $n + k - 1$ éléments $1, 2, 3, \dots, n, J_1, J_2, \dots, J_{k-1}$. Par conséquent on trouve $\bar{C}(k, n) = C(k, n + k - 1)$. On a donc:

Satz • Théorème 3.10

Kombinationen mit Wiederholung: • Combinaisons avec répétitions:

$$\bar{C}(k, n) = C(k, n + k - 1) = \binom{n + k - 1}{k}$$

Beispiel: • Exemple

Ein Abteilungsleiter hat 19 Ingenieure unter sich, von denen jeder als Projektleiter in Frage kommt. Es stehen 8 neue Projekte an, die wahrscheinlich nacheinander bearbeitet werden müssen. Wieviele Möglichkeiten bieten sich dem Abteilungsleiter, Projektleiter zu bestimmen, wenn auch in Betracht gezogen werden darf, dass im Extremfall derselbe Ingenieur allen 8 Projekten vorsteht?

• *Un chef de rayon dirige 19 ingénieurs desquels chacun est capable d'avoir la responsabilité pour un projet. 8 nouveaux projets sont à faire (en suspens), qui doivent être traités vraisemblablement l'un après l'autre. Combien de possibilités s'offrent au chef de rayon de nommer des responsables pour les projets, si on peut tirer en considération dans le cas extrême, que le même ingénieur assume (dirige) tous les 8 projets?*

Hier handelt es sich um eine Kombination mit Wiederholung. Aus 19 Ingenieuren werden 8 Projektleiter ausgewählt, wobei jeder mehrmals vorkommen darf. Es ist dann:

• *Ici, il s'agit d'une combinaison avec répétitions. Dans un ensemble de 19 ingénieurs, 8 responsables sont choisis de façon que chacun peut être nommé plusieurs fois. Il est donc:*

$$\bar{C}(8, 19) = \binom{19 + 8 - 1}{8} = \binom{26}{8} = \frac{26!}{8! \cdot (26 - 8)!} = \frac{26!}{8! \cdot 18!} = 1562275 \approx 1.56228 \cdot 10^6.$$

Neuere Darstellung siehe: • *Nouvel exposé allemand voir:*

<http://rowicus.ch/Wir/Scripts/TEIL6dCrashKursWahrschKomb.pdf>

3.4 Übungen — Exercices

Übungen finden sich in *DIYMU*, Wirz, Bibl. A18 sowie in der klassischen Schulbuchliteratur für die Gymnasialstufe — oder speziell auch in der Literatur zur Wahrscheinlichkeitsrechnung und Statistik.

• *On trouve des exercices dans DIYMU, Wirz, Bibl. A18, ainsi que dans la littérature scolaire classique pour le niveau gymnasial ou spécialement aussi dans les manuels du calcul des probabilités et statistiques.*

Kapitel 4

Wahrscheinlichkeitsrechnung — Calcul des probabilités

4.1 Einleitung — Introduction

4.1.1 Problemstellung — Problème

Gesucht ist ein Instrument, das Aussagen über betrachtete Massenereignisse erlaubt.

- *On cherche un instrument qui permet de faire des proposition concernant des événements de masse.*

Beispiele von Massenereignissen: • *Exemples d'événements de masse:*

1. Grosse Zahl von Individuen
 - *Grand de nombre d'individus*
2. Grosse Zahl von zeitlich gestaffelten Einzelercheinungen
 - *Grand nombre d'apparences individuelles temporellement échelonnées*
3. Grosse Zahl von räumlich gestaffelten Einzelercheinungen
 - *Grand nombre d'apparences individuelles échelonnées dans l'espace*
4. Grosse Zahl von Möglichkeiten
 - *Grand nombre de possibilités*
5. u.s.w. • *etc.*

Zudem suchen wir Begriffe, die bei der Vermessung solcher **Ereignismengen** nützlich sind.

- *En outre nous cherchons des notions qui sont utiles pour mesurer ces ensembles d'événements.*

4.1.2 Anwendung — Application

Wissenschaften, in denen solche Massenereignisse Gegenstand der Betrachtung sind, sind z.B. die folgenden:

- *Les sciences, dans lesquelles de tels événements de masse sont objet de la considération, sont par exemple les suivantes:*

1. Statistik • *Les statistiques*
2. Prognosenlehre • *La science de la prévision*

3. Technik der Optimierung von Verteilungen (z.B. Datenverteilung auf Datenträgern)
 - *La technique de l'optimisation de dispositions (par exemple la disposition de données sur les supports de données)*
4. Modellierung in Ökonomie, Finanzmathematik, Versicherungslehre, Technik, Medizin u.s.w.. (z.B. Sterblichkeitsvorhersagen, Börsenverhalten, Therapiechancen u.s.w.)
 - *Créer des modèles en économie, mathématiques financière, théorie d'assurance, technique, médecine etc.. (par exemple prédictions de mortalité, comportement de bourse, chances de thérapie etc.)*
5. u.s.w. • *etc.*

Das mathematische Mittel, auf denen die genannten Theorien fussen, ist die **Wahrscheinlichkeitsrechnung**.

- *Le moyen mathématique, sur lequel les théories mentionnées se fondent, est le calcul des probabilités.*

Beispiel Statistik: • **Exemple de la statistique:**

Zur blossen Beschreibung von Massenerscheinungen dient die **deskriptive** oder **beschreibende Statistik**.

- *Quant à la description simple des événements de masse, on se sert de la statistique descriptive.*

Will man hingegen Massenerscheinungen auch noch beurteilen, vergleichen oder ihr Verhalten modellieren, so benützt man die **mathematische Statistik** (beurteilende oder affirmative Statistik, Teststatistik, explorative Statistik, Tests von Modellen wie z.B. Regressionen u.s.w.). Grundlage der mathematischen Statistik ist die **Wahrscheinlichkeitsrechnung**.

Wichtige Idee: Gesucht sind mathematische Methoden, die bei einem vertretbaren Arbeitsaufwand einen verlässlichen Schluss vom Verhalten weniger Individuen (Stichprobe) auf die Grundgesamtheit (sehr viele Individuen) zulassen.

Statistik gehört auch zur Allgemeinbildung des Ingenieurs. Sie ist vor allem wichtig in der Qualitätskontrolle und im Bereich des Vertrauens in technische Komponenten (Grundlage des Vertrauens).

- *Si par contre on veut juger ou comparer des événements de masse ou bien étudier leur comportement à l'aide d'un modèle, on utilise les statistiques mathématiques, statistiques qui servent à juger ou affirmatives (statistiques de test, statistiques exploratives, tests de modèles, par exemple comme régressions etc.). Base des statistiques mathématiques est le calcul des probabilités.*

Idee importante: *On cherche des méthodes mathématiques qui admettent une conclusion sûre ou bien fiable du comportement de peu d'individus (échantillons prélevés) sur la totalité des individus (lots, beaucoup d'individus) sans dépenser trop de travail.*

Aujourd'hui les méthodes statistiques font partie de la culture générale pour l'ingénieur. Elles sont importantes pour le contrôle de la qualité et aussi, quant aux composantes techniques, dans la théorie de la fiabilité.

4.1.3 Personen — Personnages

Personen in der Entstehungsgeschichte der Wahrscheinlichkeitsrechnung:

- *Les personnages de l'histoire du calcul des probabilités:*
 - Plaise Pascal (1623 – 1662)
 - Pierre Fermat (1601 – 1665)
 - Christian Huygens (1629 – 1695)
 - Jakob Bernoulli (1654 – 1705) Def. der Wahrscheinlichkeit • *Déf. de la probabilité*
 - Abraham Moivre (1667 – 1754) Def. der Wahrscheinlichkeit • *Déf. de la probabilité*

- Pierre Simon Laplace (1749 – 1827) Def. der Wahrscheinlichkeit • *Déf. de la probabilité*
- Carl Friederich Gauss (1777 – 1855)
- Simon–Denis Poisson (1781 – 1840)
- Pafnuti Lwowitsch Tschebyschow (1821 – 1894)
- Andrei Andrejewitsch Markow (1856 – 1922)
- Andrei Nikolajewitsch Kolmogorow (1903 – 1987) Def. der Wahrscheinlichkeit • *Déf. de la probabilité*
- Alexander Jakowlewitsch Chintchin (1894 – 1959)

4.2 Zufallsexperiment, Ereignis — Expérience de hasard, événement

Begriff: Zufallsvorgang: • **Notion: Événement au hasard:**

Zwei Interpretationen sind möglich: • *Deux Interprétations sont possibles*

1. **Relativer Zufall:** Vorgang, der einen Grund hat, dessen Grund jedoch unbekannt ist.
 - **Hasard relatif:** *Processus, qui a une raison mais dont la raison est inconnue.*
2. **Absoluter Zufall:** Vorgang, der aus sich selbst heraus existiert und der prinzipiell keinen weiteren Grund hat, also von nichts abhängig ist¹.
 - **Hasard relatif absolu:** *Processus, qui n'existe que de lui-même et qui en principe n'a aucune raison, et qui ne dépend donc de rien*².

Das Problem der Existenz des absoluten Zufalls gehört in die Philosophie. Denn aus der Unkenntnis einer Sache lässt sich nicht ihre Nichtexistenz folgern.

- *Le problème de l'existence du hasard absolu est un problème de la philosophie. Car si une chose n'est pas connue, il n'est pas possible d'en conclure de son inexistence.*

Definition: • **Définition:**

Zufallsexperiment oder Zufallsbeobachtung:

• **Expérience de hasard ou observation de hasard:**

Beliebig oft wiederholbarer, nach festen Vorschriften ausgeführter Vorgang mit nicht vorherbestimmbarem Ergebnis (Ergebnismenge) und mehreren möglichen, sich gegenseitig ausschliessenden Resultaten. („Abhängig vom Zufall“.)

• *Processus qui est réalisé d'après des prescriptions fixes (ensemble d'événements), réitérable aussi souvent qu'on veut, avec un résultat non-définissable d'avance et qui a plusieurs résultats possibles qui s'excluent les uns les autres. ("Dépendant du hasard".)*

Bsp.: • **Exemple:**

Würfeln, werfen einer Münze, Kartenziehen, Grösse von Menschen messen, Milchmenge einer Kuh messen, farbige Kugeln aus einer Urne ziehen u.s.w..

- *Jouer aux dés, jeter une monnaie, tirage de cartes à jouer, mesurer la taille des gens, mesurer la quantité de lait au'une vache produit, tirer des boules colorées d'une urne etc..*

Sprechweise: • **Façon de dire:**

Das Ereignis „hängt vom Zufall ab“ := Das Ergebnis ist nicht im voraus bestimmbar.

- *L'événement "dépend du hasard" := Le résultat n'est pas prévisible.*

Bsp.: • **Exemple:** „Würfeln der Zahl 6“.

- *Jouer au dés et faire un six.*

¹Für viele Theologen ist dies die Eigenschaft Gottes.

²Pour beaucoup de théologiens c'est la qualité de Dieu.

Die Verteilungsfunktion liefert hier in der Regel keine Handhabe oder kein Merkmal, nach dem man ein Ereignis, z.B. würfeln einer 6, als „vor anderen solchen ausgezeichnet“ bewerten könnte (z.B. würfeln einer 5). Mangels besseren Wissens, aus pragmatischen Gründen also, müsste man daher von „gleichen Chancen für solche Ereignisse“ sprechen.

• *La fonction de répartition ne fournit ici normalement aucun maniement ou aucune caractéristique, d'après lequel on pourrait distinguer un événement parmi d'autres semblables. Par exemple jouer aux dés un 6 et jouer un 5. Par manque de connaissances, donc par des raisons pragmatiques, on doit conséquemment parler de mêmes chances pour de tels événements.*

Definition: • **Définition:**

Elementarergebnis: • **Résultat élémentaire:**

Mögliches, nicht weiter zerlegbares Ergebnis (ω_i) eines Zufallsexperiments.

• *Résultat d'une expérience de hasard possible (ω_i) qui ne peut plus être partagé en des résultats partiels. (Mögliche Zerlegung verändert Experiment.)*

Bsp.: • **Exemple:** Würfeln mit einem Würfel $\rightsquigarrow$ 6 Möglichkeiten: $R(1), R(2), \dots, R(6)$. Oder Würfeln mit drei Würfeln. Zerlegung: Würfeln mit einem Würfel $\rightsquigarrow$ neues, **atomares Experiment**.

• **Exemple:** *Jouer aux dés avec un dé $\rightsquigarrow$ 6 possibilités: $R(1), R(2), \dots, R(6)$. Ou jouer aux dés avec trois dés. Décomposition: Jouer aux dés avec un dé $\rightsquigarrow$ expérience nouvelle et atomique.*

Wichtige Eigenschaft: • **Qualité importante:**

Ein Elementarergebnis gehorcht der zweiwertigen Logik: Es kann eintreffen oder nicht.

• *Un résultat élémentaire obéit à la logique bivalente: Il peut se réaliser ou ne pas se réaliser. (Décomp. change exp..)*

Definition: • **Définition:**

Elementarergebnismenge: Menge aller Elementarergebnisse $\Omega = \{\omega_i \mid i = \dots\}$.

• **Ensemble des résultats élémentaires:** *Ensemble de tous les résultats élémentaires $\Omega = \{\omega_i \mid i = \dots\}$.*

Bemerkung: • **Remarque:**

Statt Ergebnismenge sagen wir auch **Universalmenge** des Experiments oder **Fundamentalmenge**.

• *Au lieu de dire ensemble des résultats, nous disons aussi ensemble universel de l'expérience ou ensemble fondamental.*

Definition: • **Définition:**

Ereignis: • **Événement:**

Teilmenge der Ergebnismenge eines Zufallsexperiments.

• *Sous-ensemble de l'ensemble des résultats d'une expérience de hasard.*

Definition: • **Définition:**

Elementarereignis zum Ergebnis ω_i : $\{\omega_i\}$

• **Événement élémentaire** lié au résultat ω_i : $\{\omega_i\}$

Ereignisse sind demnach Mengen von Elementarereignissen. • *Les événements sont donc des ensembles d'événements élémentaires.*

Definition: • **Définition:**

Ereignisraum: • **Espace d'événements:**

Menge aller Ereignisse eines Zufallsexperiments. ($= \mathcal{P}(\Omega)$, Ω = Grundmenge, Ergebnismenge, $\mathcal{P}$ = Potenzmenge.)

• *Ensemble de tous les événements ($\{\text{résultats}\}$) d'une expérience de hasard. ($= \mathcal{P}(\Omega)$, Ω = ensemble fondamental, ensemble des résultats, $\mathcal{P}$ = ensemble de parties.)*

Schreibweise: • **Façon d'écrire:**

Für Ereignisse verwenden wir Grossbuchstaben ($A, B, C, \dots, X, Y, \dots$) wie für Mengen.

• *Pour les événements, nous utilisons des lettres majuscules ($A, B, C, \dots, X, Y, \dots$) comme pour les ensembles.*

1. Beispiel: • **Exemple 1:** $A = \{\text{würfeln einer 6}\} = \{\omega_6\}$. • $A = \{\text{jouer un 6 aux dés}\} = \{\omega_6\}$.

Kurz: • *Bref:* $A = \{R(6)\}$, $\Omega = \{R(1), R(2), R(3), R(4), R(5), R(6)\} = \{\omega_1, \dots, \omega_6\}$

2. Beispiel: • **Exemple 2:** Gegeben sei ein Stab der Länge 1, der zufällig in 3 Stücke gebrochen wird (Stellen x und y , $0 < x < y < 1$). Resultat: 3 Stabstücke mit den Längen x , $y - x$, $1 - y$. $\leadsto$ Ein Resultat dieses Experimentes kann somit als Punkt im Dreieck $\{(x, y) \in ((0, 1) \times (0, 1)) \mid x < y\}$ dargestellt werden.

• *Soit donné: Bâton de longueur 1 qui soit brisé au hasard en 3 parties (places x et y , $0 < x < y < 1$). Résultat: 3 bouts de bâton de longueur x , $y - x$, $1 - y$. $\leadsto$ Un résultat de cet expérience peut donc être représenté comme point dans le triangle $\{(x, y) \in ((0, 1) \times (0, 1)) \mid x < y\}$.*

Bemerkung: • **Remarque:**

Im ersten Beispiel oben ist der **Ereignisraum endlich**, im zweiten Beispiel ist er **unendlich**.

• *Dans le premier exemple l'espace des événements ou résultats est fini, dans le deuxième exemple il est infini.*

Ein mögliches Ergebnis kann bei einem Zufallsexperiment eintreffen oder nicht. Nach dem Experiment können wir die Menge der eingetroffenen Ergebnisse E (**statistische Ergebnisse, statistische(s) Ereignis(se)**) von der Menge der nicht eingetroffenen $\Omega \setminus E$ unterscheiden. $\leadsto \omega$ ist eingetroffen $\Leftrightarrow \omega \in E$, ω ist nicht eingetroffen $\Leftrightarrow \omega \notin E$

• *Lors d'une expérience aléatoire, un résultat possible peut être réalisé ou ne pas être réalisé. Après l'expérience, nous pouvons distinguer l'ensemble des résultats E (**résultats statistiques, événement(s) statistique(s)**) qui ont été réalisés de l'ensemble des résultats $\Omega \setminus E$ qui n'ont pas été réalisés. $\leadsto \omega$ s'est réalisé $\Leftrightarrow \omega \in E$, ω ne s'est pas réalisé $\Leftrightarrow \omega \notin E$*

4.2.1 Zufallsprozesse, Häufigkeit — Processus aléatoire, fréquence

Zufallsprozesse — Processus aléatoire

Z.B. das einmalige Würfeln mit einem Würfel ist ein atomares Zufallsexperiment, das nicht weiter zerlegbar ist. Neben solchen **atomaren Zufallsexperimenten** trifft man in der Praxis aber auch komplizierte Zufallsprozesse, die z.B. aus mehreren atomaren und nacheinander oder gleichzeitig ablaufenden Zufallsexperimenten zusammengesetzt sind. Man denke z.B. an eine Urne mit farbigen Kugeln, aus der nacheinander 3 mal je eine Kugel gezogen werden soll. In einer ersten Stufe ziehen wir eine Kugel aus 6 vorhandenen, in einer 2. Stufe ziehen wir eine aus 5 vorhandenen und in einer 3. Stufe eine aus 4 vorhandenen. Ein Elementarereignis des Zufallsexperiments besteht hier aus 3 sich folgenden atomaren Ergebnissen.

• *Par exemple jouer une seule fois aux dés avec un dé est une **expérience de hasard atomique** qui n'est plus décomposable. À côté de telles expériences de hasard atomiques, on trouve dans la pratique*

aussi des processus de hasard plus compliqués qui sont composés par exemple de plusieurs expériences de hasard atomiques qui se déroulent l'une après l'autre ou simultanément. Par exemple une urne remplie de boules colorées, de laquelle on tire 3 boules l'une après l'autre. A un premier niveau, nous tirons par exemple une boule sur 6 boules disponibles, à un 2. niveau on tire une boule sur 5 boules disponibles et à un 3. niveau une boule sur 4 boules disponibles. Ici un événement élémentaire de l'expérience de hasard consiste en 3 résultats atomiques qui se suivent.

Sprechweise: • **Façon de dire:**

Bei zusammengesetzten Zufallsexperimenten nennen wir die Nummer von sich folgenden atomaren Zufallsexperimenten die **Stufe** und das zusammengesetzte Experiment daher **mehrstufig**.

• Quant aux expériences de hasard composées, nous appelons le numéro des expériences de hasard atomiques dans une suite le **niveau**. Par conséquent l'expérience totale consiste en plusieurs niveaux.

Definition: • **Définition:**

Ein mehrstufiges Zufallsexperiment nennen wir **Zufallsprozess**.

• Une expériences de hasard à plusieurs niveaux s'appelle **processus aléatoire**.

Bsp.: • **Exemple:**

Ziehen wir zuerst eine schwarze, dann eine weiße und dann wieder eine schwarze Kugel, so schreiben wir als Resultat $R(s_1, w_2, s_3)$. Dieses Zufallsexperiment ist dreistufig. Der Index bedeutet die Stufe. Die Information des Niveaus ist notwendig, da die in der Urne vorhandenen Kugeln (Voraussetzung des Experiments) mit der Stufe ändern. Ein Elementarereignis eines n -stufigen Experiments besteht demnach aus einem geordneten n -Tupel von atomaren Elementarereignissen.

• Si nous tirons d'abord une boule noire, puis une blanche et ensuite encore une fois une boule noire, nous écrivons $R(n_1, b_2, n_3)$ pour le résultat. Cette expérience de hasard consiste en trois niveaux l'index indique le niveau. L'information du niveau est nécessaire parce que les boules disponibles dans l'urne (condition de l'expérience) changent avec le niveau. Un événement élémentaire d'une expérience au niveau n consiste donc en un multiplet- n arrangé d'événements élémentaires atomiques.

Häufigkeit — Fréquence

Sei das statistische Ereignis E das Resultat eines Versuches (Versuchsreihe vom Umfang n , n Teilversuche, Prozess zum Niveau n). Das gewünschte Ereignis A treffe bei $|E| = n$ Teilversuchen k mal ein. ($\leadsto$ n mögliche und k günstige Fälle.)

• L'événement statistique E soit le résultat d'une expérience (série à n tentatives, n expériences partielles, processus du niveau n). L'événement A soit réalisé k fois à n tentatives partielles ou élémentaires. (n cas possibles et k cas favorables.)

Beispiel: 12 mal würfeln: E = Menge der realisierten und nummerierten Teilresultate. 3 mal kommt die 1. A = Teilmenge der realisierten und nummerierten 1. $\leadsto |E| = n = 12$, $k = |A| = 3$.

• **Exemple:** 12 fois jouer aux dés: E = ensemble des résultats partiels réalisés et numérotés. Le 1 vient 3 fois. A = sous-ensemble des 1 réalisés et numérotés. $\leadsto |E| = n = 12$, $k = |A| = 3$.

Definition: • **Définition:**

Absolute Häufigkeit von A : • **Fréquence absolue** de A :

$$H(A) := k$$

Relative Häufigkeit von A : • **Fréquence relative** de A :

$$h(A) := \frac{k}{n} = \frac{|A|}{|E|}$$

Bsp.: • **Exemple:**

A = ziehen eines Königs aus einem gemischten Kartenspiel mit 36 Karten bei $n = 100$ Versuchen (100

mal ziehen). Sei $k = 9 \rightsquigarrow H(A) = 9, h(A) = 0.09$. Ein Ereignis bei diesem 100-stufigen Zufallsprozess ist demnach ein geordnetes 100-Tupel mit 100 atomaren Ereignissen, welche jeweils das Resultat einer einfachen Ziehung sind.

• $A = \text{tirer un roi d'un jeu de cartes bien mélangé avec 36 cartes et 100 tentatives (tirer 100 fois)}$. Soit $k = 9 \rightsquigarrow H(A) = 9, h(A) = 0.09$. Un événement dans ce processus de hasard au niveau 100 est donc un multiplet-100 arrangé avec 100 événements atomiques, qui sont chaque fois le résultat d'un tirage simple.

Folgerung: • **Conclusion:** Absolute Häufigkeiten sind Mächtigkeiten von **eingetroffenen** Ereignissen • *Les fréquences absolues sont des puissances d'événements qui ont eu lieu.*

Bsp.: • **Exemple:**

Versuch: 12 mal würfeln. Resultat: $E = \{2_1, 4_2, 5_3, 1_4, 2_5, 6_6, 3_7, 1_8, 1_9, 5_{10}, \cdot_{11}, 3_{12}\}$. Teilmenge der Elemente „1 ist eingetroffen“: $A := \{x_i \in E \mid x_i = 1_i\} = \{1_4, 1_8, 1_9\}$. ($\cdot_{11} = \cdot =$ „missing value“.)

$\rightsquigarrow H(A) = 3, h(A) = \frac{3}{12}$.

• *Expérience: jouer 12 fois au dé. Résultat: $E = \{2_1, 4_2, 5_3, 1_4, 2_5, 6_6, 3_7, 1_8, 1_9, 5_{10}, \cdot_{11}, 3_{12}\}$. Sous-ensemble "1 a eu lieu": $A := \{x_i \in E \mid x_i = 1_i\} = \{1_4, 1_8, 1_9\}$. ($\cdot_{11} = \cdot =$ "missing value".)*

$\rightsquigarrow H(A) = 3, h(A) = \frac{3}{12}$.

Sprechweise: • **Façon de dire:**

Die betrachtete Teilmenge A von E , der unser besonderes Interesse gilt, nennen wir hier auch **Zielmenge**.

• *Le sous-ensemble A de E objet de notre intérêt particulier se dirige, s'appelle ici ensemble de but.*

Daher gilt: • *Il vaut donc:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Sei E ein eingetroffenes Ereignis eines Zufallsprozesses, A Zielmenge.

• *Soit E un événement qui a eu lieu d'un processus aléatoire, A soit un ensemble de but.*

Beh.: • **Thè.:**

1. $H(A) = |A|$ (Absolute Häufigkeit = Mächtigkeit)

• *(Fréquence absolue = puissance)*

2. $h(A) = \frac{H(A)}{H(E)} = \frac{|A|}{|E|}$ (Relative Häufigkeit)

• *(Fréquence relative)*

3. $h(E) = 1$

Unmittelbar erkennt man folgenden Sachverhalt: • *On reconnaît aussitôt:*

Satz: • **Théorème:** $\forall_A : 0 \leq h(A) \leq 1$

4.3 Ereignisalgebra — Algèbre des événements

4.3.1 Ereignisalgebra und Mengenalgebra — Algèbre des événements et algèbre des ensembles

Da Ereignisse Mengen sind, ist Ereignisalgebra nichts anderes als Mengenalgebra.

• *Comme les événements sont des ensembles, l'algèbre des événements n'est rien d'autre que l'algèbre des ensembles.*

Seien A und B Ereignisse ($\leadsto$ Mengen), die bei einem geplanten Experiment eintreffen können.

• *Soient A et B des événements ($\leadsto$ ensembles), qui peuvent se réaliser lors d'une expérience projetée.*

Definition: • **Définition:**

Summenereignis: • **Événement somme:**

$A + B$ oder $A \cup B$:= Gesamtereignis, das genau dann eintritt, wenn A oder B oder beide gemeinsam eintreffen. ($\leadsto$ Vereinigung von A und B .)

• $A + B$ ou $A \cup B$:= événement total, est réalisé exactement si A ou B ou les deux ensembles sont réalisés. ($\leadsto$ Réunion de A et B .)

Definition: • **Définition:**

Produktereignis: • **Événement produit:**

$A \cdot B$ oder $A \cap B$:= Gesamtereignis, das genau dann eintritt, wenn A und B gemeinsam eintreffen. ($\leadsto$ Durchschnitt von A und B .)

• $A \cdot B$ ou $A \cap B$:= événement total, est réalisé exactement si A et B sont réalisés. ($\leadsto$ Intersection de A et B .)

Sei dazu Ω die Menge aller Elementarergebnisse (Universalmenge, Summenereignis aller atomaren Ereignisse, trifft stets ein) und $\{\}$ die leere Menge (trifft nie ein).

• *Soit Ω l'ensemble de tous les résultats élémentaires (ensemble universel, événement somme de tous les événements atomiques, est toujours réalisé) et $\{\}$ l'ensemble vide (n'est jamais réalisé).*

Definition: • **Définition:**

Ein stets eintreffendes Ereignis (Ω) heisst **sicheres Ereignis**. Ein Ereignis, das nie eintreffen kann ($\{\}$) heisst **unmögliches Ereignis**.

• *Un événement qui est toujours réalisé (Ω) s'appelle événement sûr. Un événement qui n'est jamais réalisé ($\{\}$) s'appelle événement impossible.*

Definition: • **Définition:**

Folgeereignis: • **Événement entraîné:**

A heisst **Folgeereignis** von B , wenn das Eintreffen von B immer das von A zur Folge hat ($B \supseteq A$).

• B s'appelle **entraîné** par A , si la réalisation de A a toujours comme conséquence la réalisation de B ($B \supseteq A$).

Definition: • **Définition:**

Komplementäres Ereignis: • **Événement complémentaire:**

$\bar{A} = A^c$ heisst **Komplementäres Ereignis** von A , wenn gilt:

$A \cup \bar{A} = \Omega \wedge A \cap \bar{A} = \{\}$.

• $\bar{A} = A^c$ s'appelle **événement complémentaire** de A , si $A \cup \bar{A} = \Omega \wedge A \cap \bar{A} = \{\}$.

Definition: • Définition:

Äquivalente Ereignisse: • **Événements équivalents:**

A und B heissen **äquivalent**, wenn das Eintreffen von A immer das von B zur Folge hat und umgekehrt. $\leadsto (B \supseteq A) \wedge (A \supseteq B)$.

• A et B s'appellent **équivalents**, si la réalisation de A a toujours comme conséquence la réalisation de B et vice versa. $\leadsto (B \supseteq A) \wedge (A \supseteq B)$.

Definition: • Définition:

Ausschliessende Ereignisse: • **Événements s'excluant mutuellement:**

A und B heissen **ausschliessend**, falls A und B nie gleichzeitig eintreffen können. ($A \cap B = \{\}$.)

• A et B s'appellent **s'excluant mutuellement (incompatibles)**, si A et B ne sont jamais réalisés simultanément. ($A \cap B = \{\}$.)

1. Beispiel: • Exemple 1:

Zufallsexperiment „ein Kind kommt zur Welt“: $\leadsto$ Zwei Möglichkeiten: Ereignis $A = \{\text{Es wird ein Junge}\}$. Ereignis $B = \{\text{Es wird ein Mädchen}\}$.

• L'expérience de hasard "un enfant va naître": $\leadsto$ Deux possibilités: événement $A = \{\text{Ça va être un garçon}\}$. L'événement $B = \{\text{Ça va être une fille}\}$.

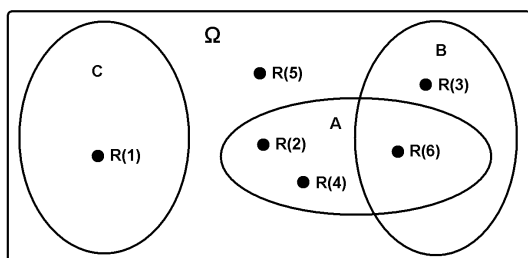
$\leadsto A \cup B = \text{sicheres Ereignis}$, $A \cap B = \text{unmögliches Ereignis}$. A, B sind ausschliessend.

• $\leadsto A \cup B = \text{événement sûr}$, $A \cap B = \text{événement impossible}$. A, B sont s'excluent mutuellement.

2. Beispiel: • Exemple 2:

$A = \{\text{würfeln einer geraden Zahl}\}$ (kurz $R(2 \vee 4 \vee 6)$), $B = \{\text{würfeln einer durch 3 teilbaren Zahl}\}$ (kurz $R(6 \vee 3)$), $C = \{\text{würfeln der Zahl 1}\}$ (kurz $R(1)$).

• $A = \{\text{jouer aux dés un nombre pair}\}$ (bref $R(2 \vee 4 \vee 6)$), $B = \{\text{jouer aux dés un nombre divisible par 3}\}$ (bref $R(6 \vee 3)$), $C = \{\text{jouer aux dés le nombre 1}\}$ (bref $R(1)$).



$$A \cup B = R(2 \vee 3 \vee 4 \vee 6)$$

$$A \cap B = R(6)$$

Bemerkung: • Remarque:

Eine Algebra ist eine Menge zusammen mit Operationen auf den Elementen. Die Potenzmenge von Ω bildet daher eine Algebra mit Operationen auf ihren Teilmengen, welche Mengen von Ereignissen sind. Die Operationen sind $\cup, \cap, \overline{\cdot}, \subset$. Wir sprechen also hier von einer **Ereignisalgebra**.

• Une algèbre est un ensemble avec opérations sur les éléments. L'ensemble de parties de Ω forme par conséquent une algèbre avec des opérations sur leurs sous-ensembles qui sont des événements. Les opérations sont $\cup, \cap, \overline{\cdot}, \subset$. Nous parlons donc ici d'une **algèbre d'événements**.

4.3.2 Boolesche Algebren — Algèbres de Boole

Boolesche Algebren sind in Wirz, Bibl. A14 (Mathematikurs für Ingenieure, Teil 4, Einführung in die Boolesche Algebra) besprochen. Wir resumieren hier nur die in diesem Kurs wichtige Anwendung bezüglich der Mengenalgebra. • *L'algèbre de Boole est traitée dans Wirz, Bibl. A14 (Mathematikurs für Ingenieure, Teil 4, Einführung in die Boolesche Algebra). Ici nous résumons seulement l'application concernant l'algèbre des ensembles importante dans ce cours.*

Sei • *Soit* $\mathcal{A} \subset \mathcal{P}(\Omega)$, $\{\} \in \mathcal{A}$. ($\mathcal{P}(\Omega)$ = Potenzmenge) • ($\mathcal{P}(\Omega)$ = ensemble de parties)

Für unsere Bedürfnisse definieren wir: • *Nous définissons pour nos besoins:*

Definition: • **Définition:** $\mathcal{A}$ heisst **boolesche Algebra** (boolesche Mengenalgebra), wenn gilt:
• $\mathcal{A}$ s'appelle **algèbre de Boole** (algèbre des ensembles de Boole), si l'on a:

1. $A \in \mathcal{A} \Rightarrow \overline{A} \in \mathcal{A}$
2. $(A \in \mathcal{A}) \wedge (B \in \mathcal{A}) \Rightarrow (A \cup B) \in \mathcal{A}$

Es gilt (De Morgan!): • *Il vaut (De Morgan!):*

$$(A \in \mathcal{A}) \wedge (B \in \mathcal{A}) \Rightarrow (\overline{A} \in \mathcal{A}) \wedge (\overline{B} \in \mathcal{A}) \Rightarrow (\overline{A \cup B}) \in \mathcal{A} \Rightarrow \overline{(\overline{A} \cap \overline{B})} = (\overline{\overline{A}} \cap \overline{\overline{B}}) = A \cap B \in \mathcal{A}$$

Konsequenz: • **Conséquence:** $(A \in \mathcal{A}) \wedge (B \in \mathcal{A}) \Rightarrow A \cap B \in \mathcal{A}$

Folgerungen:

• **Conclusions:**

1. $(A_i \in \mathcal{A}), i = 1, \dots, n \Rightarrow (\bigcup_{i=1}^n A_i) \in \mathcal{A}$
2. $(A_i \in \mathcal{A}), i = 1, \dots, n \Rightarrow (\bigcap_{i=1}^n A_i) \in \mathcal{A}$

Problem: • **Problème:** Wenn Ω endlich ist, ist $\mathcal{P}(\Omega)$ ebenfalls endlich und daher problemlos zu bilden. $\mathcal{P}(\Omega)$ ist eine Menge von Teilmengen von Ω , also eine **Menge von Ereignissen**. Wie bildet man aber $\mathcal{P}(\Omega)$, wenn Ω unendlich ist? Und gilt dann $(\bigcup_{i=1}^{\infty} A_i) \in \mathcal{A}$ immer noch?

• *Soit Ω fini. Par conséquent $\mathcal{P}(\Omega)$ est aussi fini et donc à former sans problèmes. $\mathcal{P}(\Omega)$ est un ensemble de sousensembles de Ω et par conséquent un **ensemble d'événements**. Mais comment former $\mathcal{P}(\Omega)$, si Ω est infini? Et est-ce que $(\bigcup_{i=1}^{\infty} A_i) \in \mathcal{A}$ reste toujours valable?*

Wie oft in der Mathematik umgehen wir das Problem mit einer Definition:

• *Combien souvent dans les mathématiques nous éludons le problème avec une définition:*

Definition: • **Définition:** Eine boolesche Mengenalgebra heisst **σ -Algebra**, wenn gilt:
• *Une algèbre des ensembles de Boole s'appelle σ -algèbre, s'il vaut:*
 $(A_i \in \mathcal{A}, i = 1, \dots, \infty) \Rightarrow (\bigcup_{i=1}^{\infty} A_i) \in \mathcal{A}$

Wegen De Morgan gilt dann wieder: • *A cause de De Morgan il vaut donc:*

Folgerung: • **Conclusion:** $(A_i \in \mathcal{A}, i = 1, \dots, \infty) \Rightarrow (\bigcap_{i=1}^{\infty} A_i) \in \mathcal{A}$

Bemerkung: • **Remarque:** Wenn Ω unendlich ist, ist $\mathcal{A}$ im allgemeinen eine Teilmenge von $\mathcal{P}(\Omega)$. Denn die Ereignisse in $\mathcal{A}$ sind normalerweise endliche Mengen, $\mathcal{P}(\Omega)$ besitzt hingegen in diesem Fall auch unendliche Teilmengen.
 • *Si Ω est infini, $\mathcal{A}$ est généralement un sous-ensemble de $\mathcal{P}(\Omega)$. Car les événements de $\mathcal{A}$ sont normalement des ensembles finis, $\mathcal{P}(\Omega)$ possède par contre dans ce cas aussi des ensembles infinis.*

Konsequenz: • **Conséquence:**

Zu einem Zufallsexperiment gehört von nun an immer ein Paar $(\Omega, \mathcal{A})$. (Menge der Ergebnisse Ω und Ereignisalgebra $\mathcal{A}$.)

• *A une expérience aléatoire nous associons pour le moment toujours un couple $(\Omega, \mathcal{A})$. (Ensemble des résultats Ω et algèbre des événements $(\Omega, \mathcal{A})$.)*

4.3.3 Zur Mächtigkeit und Häufigkeit — Quant à la puissance et à la fréquence

Seien nun A und B als Zielmengen Teilmengen derselben Ergebnismenge E . Aus der Mengenalgebra können wir folgern: • *Soient maintenant A et B des sous-ensembles du même ensemble de résultats E . De l'algèbre des ensembles nous pouvons conclure:*

Satz: • **Théorème:**

1. $H(A \cup B) = H(A) + H(B) - H(A \cap B)$
2. $h(A \cup B) = h(A) + h(B) - h(A \cap B)$
3. $A \cap B = \{\}$ $\Rightarrow H(A \cup B) = H(A) + H(B)$ (ausschliessende Ereignisse)
 • *(Événements incompatibles)*
4. $A \cap B = \{\}$ $\Rightarrow h(A \cup B) = h(A) + h(B)$ (ausschliessende Ereignisse)
 • *(Événements incompatibles)*
5. $\overline{A \cup B} = \overline{A} \cap \overline{B}$ (De Morgan)
6. $\overline{A \cap B} = \overline{A} \cup \overline{B}$ (De Morgan)
7. $h(\overline{A}) = 1 - h(A)$ (De Morgan)

Bsp.: • **Exemple:**

Versuch: 16 mal würfeln mit 2 Würfeln. $\leadsto$ Ergebnis:

• *Expérience: jouer 16 fois aux dés avec 2 dés.*

$\leadsto$ *Résultat:*

$$E = \{(2, 3)_1, (2, 5)_2, (1, 3)_3, (4, 6)_4, (1, 6)_5, (1, 6)_6, (2, 2)_7, (4, 6)_8, (4, 6)_9, (5, 6)_{10}, (1, 1)_{11}, (2, 6)_{12}, (2, 5)_{13}, (1, 1)_{14}, (1, 5)_{15}, (3, 4)_{16}\}$$

A : Mindestens eine Zahl ist durch 3 teilbar. • *Au moins un nombre est divisible par 3.*

B : Beide Zahlen sind gerade. • *Les deux nombres sont pairs.*

$$\leadsto H(E) = 16, H(A) = 10, H(B) = 5, H(A \cap B) = 4, H(A \cup B) = 11, H(\overline{A}) = 6, \dots$$

4.3.4 Ereignisbäume — Des arbres d'événements

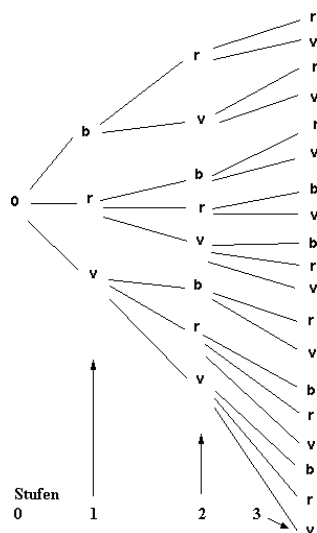
Um die Übersicht über n -stufige Zufallsprozesse zu gewinnen, ist die Darstellung der Ereignisse durch **Ereignisbäume** günstig.

- *Pour obtenir une vue d'ensemble sur un processus aléatoire du niveau total n , la représentation des événements par un arbre d'événements est favorable.*

Bsp.: • Exemple:

Gegeben sei eine Urne, die eine blaue (b), zwei rote (r) und drei violette (v) Kugeln enthält. Es sollen in einem dreistufigen Zufallsprozess nacheinander drei Kugeln gezogen werden. Die Ereignisse sollen in einem Ereignisbaum dargestellt werden. Wieviele Möglichkeiten gibt es?

- *Soit donnée une urne qui contient une boule bleu (b), deux boules rouges (r) et trois boules violetes (v). Par un processus aléatoire du niveau total trois, il faut tirer trois boules l'une après l'autre. Les événements doivent être représentés dans un arbre d'événements. Combien de possibilités est-ce qu'il y a?*



Bsp.: • Exemple:

$b - r - r$ ist eines der Ergebnisse.

Es gibt 19 verschiedene Ergebnisse (Variationen: Auswahl mit elementabhängiger Wiederholung und Anordnung).

- $b - r - r$ Est un résultat.

Il y a 19 résultats différents (Arrangements: Combinaison avec répétition dépendant de l'élément et rangement).

Wichtig: • Important:

Da die Wiederholung der ausgewählten Elemente elementabhängig ist, kann keine der 6 klassischen Formeln der Kombinatorik direkt angewandt werden, um die Anzahl der Möglichkeiten zu berechnen.

- *Comme la répétition des éléments choisis est dépendante des éléments, aucune des 6 formules classiques de l'analyse combinatoire ne peut être appliquée directement pour calculer le nombre de possibilités.*

4.4 Klassische Wahrscheinlichkeit nach Laplace — Probabilité classique d'après Laplace

4.4.1 Laplace-Experiment, Gleichwahrscheinlichkeit — Expériences de Laplace, probabilité identique

Im Folgenden studieren wir **Laplace-Experimente**. Das sind Experimente mit einer endlichen Ergebnismenge, bei denen man beobachtet, dass bei genügend grosser Anzahl Wiederholungen die Elementarereignisse alle ungefähr gleich grosse Häufigkeit aufweisen.

- *Nous allons étudier les expériences de Laplace. Ce sont des expériences avec un ensemble de résultats fini où on peut observer qu'à un nombre suffisamment grand de répétitions tous les événements élémentaires ont à peu près la même fréquence.*

Z.B. beim Werfen einer Münze hat man beobachtet • *P. ex. on a observé en jetant d'une monnaie* (Kreyszig, Bibl. A9, Statistische Methoden und ihre Anwendungen, p. 58):

Experimentator • <i>Expérimentateur</i>	Anzahl Würfe • <i>Nombre de jets</i>	Eintreffen von Kopf • <i>Face (tête) réalisée</i>	$h(R(Kopf))$ • $h(R(tête))$
Buffon	4040	2048	0.5069
Pearson	12000	6019	0.5016
Pearson	24000	12012	0.5005

Hier kann man demnach von einem Laplace-Experiment ausgehen. Man spricht auch von **statistischer Regelmässigkeit** oder von der **Stabilität der relativen Häufigkeit**. Ähnliche Erfahrungen macht man auch mit einem idealen Würfel.

• *Ici, on peut donc parler d'une expérience de Laplace. On parle aussi de la régularité statistique ou de la stabilité de la fréquence relative. On fait des expériences pareilles avec un dé idéal.*

Den Begriff des Laplace-Experiments können wir exakter fassen durch den Begriff der **Gleichwahrscheinlichkeit**:

• *L'idée de l'expérience de Laplace peut être définie plus exactement par l'idée de la probabilité égale:*

Definition: • **Définition:**

Ist bei einem Zufallsexperiment kein Ereignis vor einem andern ausgezeichnet, so heissen alle Ereignisse **gleichwahrscheinlich** oder **gleichmöglich**.

• *Si à une expérience de hasard le comportement d'aucun événement n'est distingué du comportement des autres, on appelle tous les événements de même probabilité ou de même possibilité. (Probabilité identique.)*

Folgerung: • **Conclusion:**

Bei Laplace-Experimenten sind somit die Elementarereignisse gleichwahrscheinlich.

• *Par conséquent quant à une expérience de Laplace les événements élémentaires sont de même probabilité.*

Bsp.: • **Exemple:** Laplace-Experimente: Kartenziehen, würfeln mit idealen Würfeln, ziehen von Kugeln aus einer Urne u.s.w..

• *Expériences de Laplace: Tirer des cartes dans un jeu de cartes, jouer aux dés avec des dés idéaux, tirer des boules d'une urne etc..*

4.4.2 Wahrscheinlichkeit als Mass für Gewinnchancen — Probabilité comme mesure pour gagner

Bsp.: • **Exemple:** Bei einem idealen Würfel sind alle 6 Ereignisse gleichwahrscheinlich (Laplace-Experiment). • *Quant à un dé idéal, tous les 6 événements sont de la même probabilité (expérience de Laplace).*

Gedankengang von Laplace: • **Raisonnement de Laplace:**

Problem: In einem Würfelspiel mit einem Würfel gewinnt ein Spieler, wenn er eine 3 oder eine 6 würfelt. Sonst verliert er. Was sind seine Gewinnchancen?

• *Problème: Dans un jeu de dés avec un dé, un joueur gagne, s'il fait un 3 ou un 6. Sinon il perd. Quels sont ses chances de gagner?*

Beim genannten Spiel gibt es $m = 6$ mögliche Ereignisse (**mögliche Fälle**, Ereignismenge $\Omega = \{1, 2, 3, 4, 5, 6\}$, kurz geschrieben). Für den Gewinn günstig sind jedoch nur $g = 2$ Ereignisse

(**günstige Fälle**, Ereignismenge $A = \{3, 6\}$). $m - g = 4$ Fälle sind demnach ungünstig für den Gewinn (Zahlen 1, 2, 4, 5).

• *Au jeu mentionné, il y a $m = 6$ événements possibles (**cas possibles**, ensemble d'événements $\Omega = \{1, 2, 3, 4, 5, 6\}$, écrit brièvement). Pour gagner cependant il n'y a que $g = 2$ événements (**cas favorables**). $m - g = 4$ cas sont donc défavorables pour gagner (nombres 1, 2, 4, 5), ensemble d'événements $A = \{3, 6\}$.*

$\leadsto g = H(A) = m(A), m = H(\Omega) = m(\Omega)$

Da alle Ereignisse gleichwahrscheinlich sind, kann als Mass P für die Gewinnaussicht die relative Häufigkeit genommen werden:

• *Comme tous les événements sont de la même probabilité, on peut prendre comme mesure P pour la chance de gagner la fréquence relative:*

$$P = \frac{g}{m} = \frac{H(A)}{H(\Omega)}$$

Dieses Mass für die Gewinnchance bei gleichwahrscheinlichen Elementarereignissen definiert man nun klassisch als **Wahrscheinlichkeit** für das Eintreffen von A

• *Classiquement on définit cette mesure pour la chance de gain (aux événements élémentaires de même probabilité) comme **probabilité** de réaliser l'événement A :*

Definition: • **Définition:**

Wahrscheinlichkeit (nach Laplace) $P(A)$ des Ereignisses A bei einem Experiment mit gleichwahrscheinlichen Elementarereignissen:

• **Probabilité (d'après Laplace)** $P(A)$ de l'événement A lors d'une expérience avec des événements élémentaires de même probabilité:

$$P(A) = \frac{g}{m} = \frac{H(A)}{H(\Omega)}$$

Bsp.: • **Exemple:** Zufallsprozess: Würfeln mit 3 Würfeln in 2 Versuchen (2-stufig). Gewonnen hat, wer mindestens einmal 2 6-er würfelt (Ereignis A). $\leadsto P(A) = ?$

• *Processus aléatoire: Jouer aux dés avec 3 dés et 2 essais (niveau total 2). Gagner signifie: Faire au moins une fois (donc par niveau) 2 fois un 6 (événement A). $\leadsto P(A) = ?$*

Lösung: • **Solution:**

Wir gliedern das Problem erst in Teilprobleme auf, deren Resultate wir dann anschliessend zusammenfügen.

• *D'abord nous décomposons le problème en problèmes partiels, dont nous joignons finalement les résultats.*

1. Das Ereignis A_1 sei definiert durch den Erfolg im 1. Versuch (Stufe 1), wobei das Resultat im 2. Versuch (Stufe 2) beliebig sein kann. Statt Würfel betrachten wir die Plätze P , auf die die Resultate notiert werden. Es gibt 3 Möglichkeiten, je mit zwei Würfeln eine 6 zu würfeln und mit dem jeweils 3. Würfel keine 6 zu würfeln ($E_{1,1}, E_{2,1}, E_{3,1}$). Dazu gibt es noch eine Möglichkeit, mit allen 3 Würfeln eine 6 zu würfeln ($E_{4,1}$).

$g_{j,k}$ = Anzahl günstige Möglichkeiten beim Ereignis $E_{j,1}$ am Platz Nummer k .

• *L'événement A_1 soit défini par le succès dans le 1-er essai (niveau 1), avec la restriction que le résultat du 2-ème essai (niveau 2) peut être quelconque. Au lieu de considérer les dés, nous considérons les places P sur lesquelles les résultats sont notés. Il y a 3 possibilités de faire un résultat égal 6 avec deux dés et inégal 6 avec le 3-ème dé restant ($E_{1,1}, E_{2,1}, E_{3,1}$). En plus il y a encore la possibilité de faire un 6 avec tous les 3 dés ($E_{4,1}$).*

$g_{j,k}$ = Nombre de possibilités favorables à l'événement $E_{j,1}$ à la place numéro k .

Ereign. • Evén.	Stufe 1 • Niveau 1			Stufe 2 • Niveau 2		
	$P = 1$	$P = 2$	$P = 3$	$P = 4$	$P = 5$	$P = 6$
$E_{1,1}$	$g_{1,1} = 1$	$g_{1,2} = 1$	$\neg 6 : g_{1,3} = 5$	$g_{1,4} = 6$	$g_{1,5} = 6$	$g_{1,6} = 6$
$E_{2,1}$	$g_{2,1} = 1$	$\neg 6 : g_{2,2} = 5$	$g_{2,3} = 1$	$g_{2,4} = 6$	$g_{2,5} = 6$	$g_{2,6} = 6$
$E_{3,1}$	$\neg 6 : g_{3,1} = 5$	$g_{3,2} = 1$	$g_{3,3} = 1$	$g_{3,4} = 6$	$g_{3,5} = 6$	$g_{3,6} = 6$
$E_{4,1}$	$g_{4,1} = 1$	$g_{4,2} = 1$	$g_{4,3} = 1$	$g_{3,4} = 6$	$g_{3,5} = 6$	$g_{3,6} = 6$

A_1 liegt vor, wenn E_1 oder E_2 oder E_3 oder E_4 vorliegen. • A_1 est réalisé, si E_1 ou E_2 ou E_3 ou E_4 est réalisé.

$A_1 = E_1 \cup E_2 \cup E_3 \cup E_4 \Rightarrow m(A) = g = 5 \cdot 6 \cdot 6 \cdot 6 + 5 \cdot 6 \cdot 6 \cdot 6 + 5 \cdot 6 \cdot 6 \cdot 6 + 1 \cdot 6 \cdot 6 \cdot 6 = (3 \cdot 5 + 1) \cdot 6^3$.
Da auf jedem Platz 6 Zahlen möglich sind, gilt: • Comme sur chaque place 6 nombres sont possibles, il vaut: $m = 6 \cdot 6 \cdot 6 \cdot 6 \cdot 6 \cdot 6 = 6^6$.

$$\Rightarrow P(A_1) = h(A_1) = \frac{m(A_1)}{E} = \frac{(3 \cdot 5 + 1) \cdot 6^3}{6^6} = \frac{16}{6^3} = \frac{8}{3 \cdot 6^2}$$

2. Das Ereignis A_2 sei definiert durch den Erfolg im 2. Versuch, wobei das Resultat im 1. Versuch beliebig sein kann. Aus Symmetriegründen erhalten wir dasselbe Resultat wie im 1. Versuch.

• L'événement A_2 soit défini par le succès dans le 2-ème essai, avec la restriction que le résultat du 1-er essai peut être quelconque. Pour des raisons de symétrie, nous obtenons le même résultat comme dans le 1-er essai.

$$\Rightarrow P(A_2) = P(A_1) = h(A_2) = h(A_1) = \frac{8}{3 \cdot 6^2}$$

3. Sei • Soit $A_3 = A_1 \cap A_2$. A_3 bedeutet, dass der Erfolg sich im 1. und im 2. Versuch (Stufe) einstellen. • A_3 signifie qu'on a du succès au 1-er et au 2-ème essai (niveau).

Für den Erfolg auf den Plätzen 1, 2, 3 (1. Stufe) gibt es $3 \cdot 5 + 1 = 16$ Möglichkeiten. Zu jeder dieser 16 Möglichkeiten gibt es wieder 16 Möglichkeiten für den Erfolg auf den Plätzen 4, 5, 6 (2. Stufe). Somit gibt es für den Erfolg auf beiden Stufen (A_3) $16 \cdot 16$ Möglichkeiten.

• Pour le succès sur les places 1, 2, 3 (1-er niveau) il y a $3 \cdot 5 + 1 = 16$ possibilités. Pour chacune de ces 16 possibilités il existe encore 16 possibilités pour le succès sur les places 4, 5, 6 (2-ème niveau). Par conséquent pour le succès sur deux niveaux (A_3) il y a $16 \cdot 16$ possibilités.

$$\Rightarrow P(A_3) = P(A_1 \cap A_2) = h(A_1 \cap A_2) = \frac{m(A_3)}{E} = \frac{16 \cdot 16}{6^6} = \frac{8 \cdot 8}{3^2 \cdot 6^4}$$

4. $P(A) = P(A_1 \cup A_2) = h(A_1 \cup A_2) = h(A_1) + h(A_2) - h(A_1 \cap A_2)$
 $= \frac{8}{3 \cdot 6^2} + \frac{8}{3 \cdot 6^2} - \frac{8 \cdot 8}{3^2 \cdot 6^4} = \frac{2 \cdot 8}{3 \cdot 6^2} + \frac{8 \cdot 8}{3^2 \cdot 6^4} = \frac{47}{324} \approx 0.145062$

Achtung: • Attention:

Der Begriff der Gleichwahrscheinlichkeit versagt jedoch in der Praxis meistens: Würfel sind nicht exakt symmetrisch und homogen, die Psychologie des Werfers spielt hinein, und bei vielen statistischen Vorgängen liegen überhaupt keine Laplace-Experimente vor. Man denke z.B. nur an das Wetter! Daher ist ein erweiterter Wahrscheinlichkeitsbegriff notwendig. Wir führen daher die **Wahrscheinlichkeit axiomatisch** ein nach **Kolmogoroff**.

• L'idée de la probabilité identique des événements élémentaires manque normalement dans la pratique: Les dés ne sont pas exactement symétriques et homogènes, la psychologie du lanceur joue un rôle, et beaucoup de processus statistiques ne sont pas du tout des expériences de Laplace. Qu'on pense par exemple au temps! Par conséquent une notion de probabilité dans un sens plus large est nécessaire. Par conséquent nous introduisons la notion de la **probabilité axiomatique** d'après **Kolmogoroff**.

4.5 Axiomatischer Wahrscheinlichkeitsbegriff — Probabilité axiomatique

4.5.1 Begriff, Axiome, Folgerungen — Notion, axiomes, conclusions

Ausgangspunkt ist die Erfahrung: Ergebnisse von Zufallsexperimenten unterliegen auf die Dauer

gewissen Gesetzmässigkeiten. Nach Descartes ist es vernünftig, das einfachste und plausibelste Modell des beobachteten Vorgangs (Naturvorgangs) als Gesetzmässigkeit zu akzeptieren, sofern nicht weitere Beobachtungen ein verändertes Modell erzwingen. Ein Modell in diesem Sinne ist jedoch keine kausale Naturerklärung sondern nur eine Naturbeschreibung.

• **Notre point de départ est l'expérience:** *Les résultats d'expériences de hasard suivent à la longue certaines lois. D'après Descartes, il est raisonnable d'accepter comme loi de la nature le modèle plus simple vraisemblable de l'événement observé (phénomène de nature), si d'autres observations ne nous forcent pas de changer le modèle. Un modèle dans ce sens n'est cependant pas une explication causale des phénomènes de la nature, il est seulement une description de la nature.*

Im Folgenden wollen wir uns auf Experimente oder Vorgänge einschränken, bei denen sich die relativen Häufigkeiten mit der grösser werdenden Anzahl Versuchen statistisch immer stabiler verhalten. Dabei muss jedoch keine Gleichwahrscheinlichkeit der Elementarereignisse vorausgesetzt werden. Zufällige Unregelmässigkeiten sind nur für die Elementarereignisse, nicht aber für die von der Versuchsnummer abhängige relative Häufigkeit erlaubt, wenn die Versuchsnummer gross ist.

• *Dans ce qui suit, nous voulons nous restreindre à des expériences ou des processus qui montrent des fréquences relatives qui sont par le nombre croissant statistiquement stables, si le nombre des essais est assez grand. Pourtant il n'est pas nécessaire de présupposer que les événements élémentaires sont de même probabilité. Des irrégularités accidentelles sont permises seulement pour les événements élémentaires, mais non pas pour la fréquence relative dépendante du numéro de l'expérience, si le numéro de l'expérience est assez grand.*

Grundannahme: • Hypothèse fondamentale:

Auf der Grundlage der empirischen Beobachtungen lässt sich folgende **Annahme** vertreten: Sei n die Anzahl Versuche in einem Experiment unter gleichbleibenden Bedingungen. Dann ist $\lim_{n \rightarrow \infty} h_n(A) = h_\infty$ eine mit praktischer Gewissheit angebbare reelle Zahl $\in [0, 1]$.

Achtung! Da es in der Realität der Zeit unendlich viele Versuche nie geben kann, ist diese Annahme praktisch nicht überprüfbar. Sie ist rein theoretischer Natur. An die Stelle der praktischen Gewissheit tritt hier der praktische Erfolg mit der Annahme. Darauf beruhende Statistik ist demnach in der Grundlage keine deduktive, sondern eine induktive (experimentelle) Wissenschaft.

• *Sur la base des observations empiriques, on peut supposer l'hypothèse suivante: Soit n le nombre d'essais dans une expérience à des conditions invariables. Donc $\lim_{n \rightarrow \infty} h_n(A) = h_\infty$ est un nombre réel $\in [0, 1]$ qui peut être indiqué avec une grande certitude pratique.*

Attention: *Comme dans la réalité du temps un nombre fini d'essais ne peut jamais exister, cette supposition ne peut jamais être contrôlée dans la pratique. Sa nature est purement théorique. Ici, le succès pratique prend la place de la certitude pratique de l'hypothèse. Les statistiques reposant sur ce fait ne sont donc, dans leurs bases, pas des sciences déductives, mais inductives (expérimentales).*

Bemerkung: • Remarque:

Um h_∞ zu gewinnen, ist man auf Daten, d.h. auf Erfahrung angewiesen.

• *Pour obtenir h_∞ , on a besoin de données, ç.v.d. de l'expérience.*

Die nun verwendete Wahrscheinlichkeit beruht auf einer hypothetischen Interpretation, die wir in einer Definition fassen. Sie ist wieder als Mass für ein Ereignis aufzufassen.

• *La probabilité qui va être utilisée maintenant est fondée sur une interprétation hypothétique que nous résumons dans une définition. Elle est de nouveau une sorte de mesure d'un événement.*

Definition: • **Définition:** **Wahrscheinlichkeit** oder **Wahrscheinlichkeitsmass** (statistische Definition): • **Probabilité ou mesure de probabilité** (définition statistique):

$$P : \mathcal{A} \mapsto [0, 1] \subset \mathbb{R}$$

$$A \mapsto P(A) := h_\infty = \lim_{n \rightarrow \infty} \frac{H(A)}{H(\Omega)}$$

Konsequenz: • **Conséquence:** P ist eine Funktion auf $\mathcal{A}$. • P est une fonction sur $\mathcal{A}$.

Bemerkung: • **Remarque:**

Da es keine Methode geben kann um in der praktischen Realität h_∞ zu gewinnen, setzen wir $P(A) := h_\infty \approx h_n$, $n \gg 1$ (n sehr gross).

• Comme dans la réalité pratique il ne existe pas de méthode pour obtenir h_∞ , nous posons $P(A) := h_\infty \approx h_n$, $n \ll 1$ (n très grand).

$\leadsto$ **Problem:** • **Problème:** Stabilität von h_n mit n . • Stabilité de h_n avec n .

Da Ω in der Praxis oft nicht bekannt ist, schreiben wir nun S statt Ω .

• Comme Ω est souvent inconnu dans la pratique, nous écrivons S à la place de Ω .

Axiome: • **Axiomes:**

1. $P(A) \in [0, 1]$ eindeutig bestimmt • univoque
2. (a) $A = S \Rightarrow P(A) = 1$ (sicheres Ereignis) • (événement sûr)
- (b) A und B äquivalente Ereignisse • A et B événements équivalents $\Rightarrow P(A) = P(B)$
- (c) A und B ausschliessende Ereignisse: • A et B événements incompatibles:
 $A \cap B = \{\} \Rightarrow P(A \cup B) = P(A) + P(B)$

Definition: • **Définition:** $\bar{A} = S \setminus A$ heisst **komplementäres Ereignis** von A .
 • s'appelle **événement complémentaire** de A .

Bemerkung: • **Remarque:**

1. $\Rightarrow P(A \cup B) = P(A) + P(B) - P(A \cap B)$ wegen • à cause de $A \cup B = (A \setminus B) \cup (B \setminus A) \cup (A \cap B)$ (ausschliessend!)
 • (incompatibles)
 $\leadsto P(A \cup B) \leq P(A) + P(B)$
2. Sei • Soit $A \subseteq B \Rightarrow B = A \cup (B \setminus A)$
 $\Rightarrow P(B) = P(A) + P(B \setminus A) \Rightarrow P(B) \geq P(A)$
3. $A \cap \bar{A} = \{\} \Rightarrow 1 = P(S) = P(A \cup \bar{A}) = P(A) + P(\bar{A})$
 $\Rightarrow P(A) = 1 - P(\bar{A})$
4. $P(S) = P(\bar{\bar{S}}) = 1 - P(\bar{S}) = 1 - P(\{\}) = 1 - 0 = 1$

Nachstehende Folgerungen sind nun unmittelbar einsichtig:

• *Les conclusions suivantes sont maintenant directement compréhensibles:*

Folgerungen:

• **Conclusions:**

1. $P(A) = 1 - P(\bar{A})$
2. $A_1, A_2, \dots, A_n$ ausschliessend • *incompatible*
 $\Rightarrow P(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n P(A_i)$
3. A unmöglich • *impossible* ($A = \bar{S} = \{\}$) $\Rightarrow P(A) = 0$
4. A, B beliebig • *quelconque* $\Rightarrow P(A \cup B) = P(A) + P(B) - P(A \cap B)$
5. A, B beliebig • *quelconque* $\Rightarrow P(A \cup B) \leq P(A) + P(B)$
6. $A \subseteq B \Rightarrow P(A) \leq P(B)$
7. $P(S) = 1$

Achtung: • **Attention:**

1. $0 = P(A) = h_\infty(A) \approx h_n(A)$, $n \gg 1$
 $\leadsto A$ trifft praktisch nie ein, was aber nicht bedeutet, dass A unmöglich ist!
 $\leadsto$ • *A n'est pratiquement jamais réalisé ce qui ne signifie pas que A soit impossible!*
2. $1 = P(A) = h_\infty(A) \approx h_n(A)$, $n \gg 1$
 $\leadsto A$ trifft praktisch immer ein, was aber nicht bedeutet, dass A absolut sicher ist!
 $\leadsto$ • *A est pratiquement toujours réalisé ce qui ne signifie pas que A soit absolument sûr!*

4.5.2 Der Begriff Wahrscheinlichkeitsraum — La notion espace de probabilité

Wir betrachten eine diskrete Ereignismenge $\Omega = \{\omega_1, \omega_2, \omega_3, \dots\}$ eines Zufallsexperiments. Für die Elementarereignisse ω_i sei $P(\omega_i) = p_i \geq 0$. Für $i \neq k$ muss $\omega_i \cap \omega_k = \{\}$ gelten, sonst wären die Elementarereignisse zerlegbar ($\omega_i = (\omega_i \cap \omega_k) \cup \overline{(\omega_i \cap \omega_k)}$). Daher folgt aus den Axiomen: $\sum_{i=1}^{\infty} p_i = 1$.

Da sich jedes Ereignis aus Elementarereignissen zusammensetzt ($A = \bigcup_{i=1}^{\infty} \omega_i$), gilt für die Wahrscheinlichkeit von A : $P(A) = \sum_{i=1}^{\infty} p_i$. Damit ist jedem Ereignis $A_k = \bigcup_{i=1}^{\infty} \omega_i$ eindeutig eine Wahrscheinlichkeit

$P(A_k) = \sum_{k=1}^{\infty} p_k$ zugeordnet. Diese Zuordnung $A_k \mapsto P(A_k)$ ist eine Funktion. Diese Funktion macht den Ereignisraum zu einem „**Wahrscheinlichkeitsraum**“.

• *Nous considérons un ensemble d'événements discret $\Omega = \{\omega_1, \omega_2, \omega_3, \dots\}$ lié à une expérience de hasard. Soit $P(\omega_i) = p_i \geq 0$ pour les événements élémentaires ω_i . Pour $i \neq k$ il doit être $\omega_i \cap \omega_k = \{\}$, sinon les événements élémentaires seraient démontables ($\omega_i = (\omega_i \cap \omega_k) \cup \overline{(\omega_i \cap \omega_k)}$). Par conséquent il suit des axiomes: $\sum_{i=1}^{\infty} p_i = 1$. Comme chaque événement se compose d'événements élémentaires*

($A = \bigcup_{i=1}^{\infty} \omega_i$), il vaut pour la probabilité de A : $P(A) = \sum_{i=1}^{\infty} p_i$. Avec cela, à chaque événement $A_k = \bigcup_{i=1}^{\infty} \omega_i$

il est adjoint clairement une probabilité $P(A_k) = \sum_{k=1}^{\infty} p_k$. Cette correspondance (assignation) est une fonction $A_k \mapsto P(A_k)$. Cette fonction transforme l'espace d'événements en un espace de probabilité.

Definition: • **Définition:**

Ein **Wahrscheinlichkeitsraum** ist gegeben durch einen Ereignisraum Ω , eine Ereignisalgebra und ein Wahrscheinlichkeitsmass.

• *Un espace de probabilité est donné par un espace d'événements Ω , une algèbre d'événements et une mesure de probabilité.*

Man sieht sofort, dass im Falle eines endlichen Ereignisraumes $\Omega = \{\omega_1, \omega_2, \dots, \omega_n\}$ mit $P(\omega_1) = P(\omega_2) = \dots = P(\omega_n)$ die Laplace-Wahrscheinlichkeit herauskommt.

• *On voit tout de suite que dans le cas de $\Omega = \{\omega_1, \omega_2, \dots, \omega_n\}$ avec $P(\omega_1) = P(\omega_2) = \dots = P(\omega_n)$ on obtient la probabilité de Laplace.*

Sprechweise: • **Façon de dire:**

Im Falle der Laplace-Wahrscheinlichkeit reden wir von einem **Laplace-Wahrscheinlichkeitsraum**.

• *Dans le cas de la probabilité de Laplace, on parle d'un espace de probabilité de Laplace.*

4.5.3 Bedingte Wahrscheinlichkeit — Probabilité conditionnelle

Der Begriff — **La notion****Bsp.:** • **Exemple:**

Seien A und B die folgenden Ereignisse: • *Soient A et B les ensembles suivants:*

$A = \{R(\text{Kunde ist weiblich} \mid \text{Client est féminin})\}$

$B = \{R(\text{Kunde kauft unser Produkt} \mid \text{Client achète notre produit})\}$

$\leadsto B|A = \{R(\text{Weiblicher Kunde kauft unser Produkt} \mid \text{Client féminin achète notre produit})\}$

Sprechweise: • **Façon d'exprimer:**

$B|A = \leadsto$ Ereignis B unter der Voraussetzung von Ereignis A , B abhängig von A .

• *Événement B sous la condition de l'événement A , B dépendant de A .*

Bemerkung: • **Remarque:**

$B|A \leadsto$ Kunde ist weiblich **und** Kunde kauft $\leadsto$ Schnittmenge!

• $B|A \leadsto$ Client est féminin **et** client achète $\leadsto$ intersection des ensembles!

Problem: • **Problème:** Wir fragen nun nach der Wahrscheinlichkeit, dass der Kunde unser Produkt kauft, unter der Voraussetzung, dass der Kunde weiblich ist.

• *Nous cherchons maintenant la probabilité que le client achète notre produit, à la condition que le client soit féminin.*

Definition: • **Définition:**

$P(B|A) :=$ Wahrscheinlichkeit für das Eintreffen von B unter der Voraussetzung, dass A eingetroffen ist $\leadsto$ **bedingte Wahrscheinlichkeit**.

• $P(B|A) :=$ probabilité pour la réalisation de B à la condition que A ait été réalisé $\leadsto$ **probabilité à condition**.

Zur Berechnung von $P(B|A)$ betrachten wir zuerst den Fall einer endlichen Ereignismenge Ω , $A \neq \{\}$:

• Comment calculer $P(B|A)$? Nous considérons d'abord un ensemble d'événements Ω fini, $A \neq \{\}$:

Seien • *Soient* $P(A) := \frac{H(A)}{H(\Omega)} = \frac{a}{n}$, $P(B) := \frac{H(B)}{H(\Omega)} = \frac{b}{n}$, $P(A \cap B) := \frac{H(A \cap B)}{H(\Omega)} = \frac{s}{n}$

Die Ergebnisse von $A \cap B$ treten in s Fällen ein. Dies geschieht unter der Bedingung, dass die Ergebnisse von A eingetroffen sind, was in a Fällen geschieht. A ist somit die Menge der sicheren Ereignisse

bezüglich dieses Problems, d.h. Universalmenge für das Ereignis $A \cap B$. Daher gilt:

• *Les résultats de $A \cap B$ sont réalisés dans s cas. Cela arrive à condition que les résultats de A soient réalisés ce qui arrive dans a cas. A est par conséquent l'ensemble des événements sûrs par rapport à ce problème, c.-à.-d. l'ensemble universel pour l'événement $A \cap B$. Par conséquent vaut:*

$$P(B|A) = \frac{H(A \cap B)}{H(A)} = \frac{s}{a} = \frac{\frac{s}{n}}{\frac{a}{n}} = \frac{P(A \cap B)}{P(A)}$$

Wegen $A \cap B = B \cap A$ ist obiges Gesetz symmetrisch: • *A cause de $A \cap B = B \cap A$, la loi qu'on vient de déduire est symétrique:* $P(A|B) = \frac{P(B \cap A)}{P(B)} = \frac{P(A \cap B)}{P(B)} \leadsto$

Satz: • **Théorème:**

$$\begin{aligned} 1. \quad P(B|A) &= \frac{P(A \cap B)}{P(A)} & P(A|B) &= \frac{P(B \cap A)}{P(B)} \\ 2. \quad H(A) &= H(B) \Rightarrow P(A|B) = P(B|A) \end{aligned}$$

Konsequenz: • **Conséquence:** $P(A \cap B) = P(B|A) \cdot P(A) = P(A|B) \cdot P(B)$
(Multiplikationssatz) • (Théorème de la multiplication)

Z.B. für $A = \{\}$ ist der Sachverhalt auch richtig. • *P.ex. pour $A = \{\}$ cette formule reste correcte.*
($\Rightarrow A \cap B = \{\} \Rightarrow P(A \cap B) = 0, P(A) = 0, P(B|A) = 0$)

$P(A \cap B) = P(B|A) \cdot P(A)$ lässt sich iterieren: • *le laisse itérer:*

Korollar: • **Corollaire:**

$$P(A_1 \cap A_2 \cap \dots \cap A_n) = P(A_1) \cdot P(A_2|A_1) \cdot P(A_3|A_1 \cap A_2) \cdot \dots \cdot P(A_n|A_1 \cap A_2 \cap \dots \cap A_{n-1})$$

1. Beispiel: • **Exemple 1:**

Seien • *Soient* $|\Omega| = 100, |A \cup B| = 80, |A| = 50, |B| = 45$

Es gilt: • *Il vaut:* $|A \cup B| = |A| + |B| - |A \cap B|$

$$\Rightarrow |A \cap B| = 50 + 45 - 80 = 15, \quad P(B|A) = \frac{P(A \cap B)}{P(A)} = \frac{15/100}{50/100} = \frac{15}{50} = 0.3$$

2. Beispiel: • **Exemple 2:**

Gegeben: Schachtel mit 3 grünen und 7 roten Kugeln ($\leadsto$ total 10). Wahrscheinlichkeit, dass bei 2 Zügen ohne zurücklegen nur rote Kugeln kommen?

• *Donné: Boîte avec 3 boules vertes et 7 boules rouges ($\leadsto$ en tout 10). Probabilité qu'on n'obtient que des boules rouges en tirant 2 fois sans remettre les boules tirés?*

Sei $r_1 =$ „rot im 1. Zug“, $r_2 =$ „rot im 2. Zug“. • *Soit $r_1 =$ „rouge au 1er tirage“, $r_2 =$ „rouge au 2ème tirage“.*

$$\leadsto P(\{r_1\}) = \frac{7}{10}, \quad P(\{r_2\}|\{r_1\}) = \frac{6}{9} \Rightarrow P(\{r_1\} \cap \{r_2\}) = P(\{r_2\}|\{r_1\}) \cdot P(\{r_1\}) = \frac{6}{9} \cdot \frac{7}{10} = \frac{42}{90} \approx 0.47$$

3. Beispiel: • **Exemple 3:**

Lotto: Wahrscheinlichkeit für 6 richtige aus 49? • *Lotto: Probabilité pour 6 nombres corrects de 49?*

1. Methode: Kombinationen • Méthode 1: Combinaisons

$$P = \frac{1}{\binom{49}{6}} = \frac{(49-6)! \cdot 6!}{49!} = \frac{6!}{49 \cdot 48 \cdot 47 \cdot 46 \cdot 45 \cdot 44} = \frac{1}{13983816} \approx 7.15112 \cdot 10^{-8}$$

2. Methode: Abhängige Ereignisse • Méthode 2: Événements dépendants

$A_1 \rightsquigarrow$ 1. Zahl ziehen: • *Tirage du 1er nombre*: $P(A_1) = \frac{6}{49}$

$A_2 \rightsquigarrow$ 2. Zahl ziehen unter der Voraussetzung, dass die 1. Zahl gezogen ist: • *Tirage du 2ème nombre à la condition que le 1er nombre soit tiré*: $P(A_2|A_1) = \frac{5}{48}$

$A_3 \rightsquigarrow$ 3. Zahl ziehen unter der Voraussetzung, dass die 1. und die 2. Zahl gezogen sind: • *Tirage du 3ème nombre à la condition que le 1er et le 2ème nombre soient tirés*:

$$P(A_3|A_1 \cap A_2) = \frac{4}{47}$$

$\vdots$

$$\rightsquigarrow P(A_1 \cap A_2 \cap \dots \cap A_6) = P(A_1) \cdot P(A_2|A_1) \cdot P(A_3|A_1 \cap A_2) \dots P(A_n|A_1 \cap A_2 \cap \dots \cap A_5) \\ = \frac{6}{49} \cdot \frac{5}{48} \cdot \frac{4}{47} \cdot \dots \cdot \frac{1}{44} = \frac{1}{13983816} \approx 7.15112 \cdot 10^{-8}$$

Unabhängige Ereignisse — Événements indépendants

Falls die Ergebnisse oder Elemente in der Menge $B|A$ dieselben sind wie diejenigen in der Menge B und somit $B|A = B$ gilt, hat A keinen Einfluss darauf, welche Elemente von $B|A$ bei der Auswahl aus B weggelassen werden müssen. A verändert also B nicht. Daher nennen wir B und A unabhängig. Dann wird $P(B|A) = P(B)$ und daher

$$P(A \cap B) = P(B|A) \cdot P(A) = P(A|B) \cdot P(B) = P(B) \cdot P(A) \Rightarrow P(A|B) = P(A).$$

Man kontrolliert sofort nach, dass das Ergebnis auch richtig ist für $A = \{\}$ oder $B = \{\}$.

• *Si les résultats ou les éléments dans l'ensemble $B|A$ sont les mêmes que ceux dans l'ensemble B et par conséquent s'il vaut $B|A = B$, l'ensemble A n'influence pas quels éléments de $B|A$ doivent être omis au choix dans B . Donc A ne change pas B . Par conséquent nous appelons B et A indépendants. On obtient par conséquence $P(B|A) = P(B)$ et donc*

$$P(A \cap B) = P(B|A) \cdot P(A) = P(A|B) \cdot P(B) = P(B) \cdot P(A) \Rightarrow P(A|B) = P(A).$$

On contrôle tout de suite que le résultat soit aussi correct pour $A = \{\}$ ou $B = \{\}$.

Definition: • Définition: B heisst von A **unabhängig** oder **stochastisch unabhängig**, wenn gilt: $B|A = B$.

• B s'appelle **indépendant** de A ou **indépendant de façon stochastique (aléatoire)** s'il vaut: $B|A = B$.

Satz: • Théorème: **Vor.: • Hyp.:**

Sei B von A unabhängig. • *Soit B indépendant de A .*

Beh.: • Thè.:

$$1. P(B|A) = P(B)$$

$$2. P(A \cap B) = P(A) \cdot P(B)$$

3. A von B unabhängig. • *A indépendant de B .*

1. Beispiel: • Exemple 1:

Gegeben: Kartenspiel mit 36 Karten. Ziehe 2 mal eine Karte und lege sie zurück. Wahrscheinlichkeit, dass 2 Asse kommen?

• *Donné: Jeux de cartes avec 36 cartes. Tirer 2 fois une carte et les remettre dans le jeu. Probabilité qu'on tire 2 fois un as?*

Lösung: • Solution:

a_1 : As beim 1. Zug. • *As au 1er tirage.*

a_2 : As beim 2. Zug. • *As au 2er tirage.*

Die beiden Ereignisse $A_1 = \{a_1\}$ und $A_2 = \{a_2\}$ sind unabhängig. $\leadsto P(A_1) = P(A_2)$

• *Les deux événements $A_1 = \{a_1\}$ et $A_2 = \{a_2\}$ sont indépendants. $\leadsto P(A_1) = P(A_2)$*

$$\leadsto P(A_1) = \frac{4}{36}, \quad P(A_1 \cap A_2) = P(A_1) \cdot P(A_2) = \frac{4}{36} \cdot \frac{4}{36} = \frac{1}{81}$$

Benutzung des Ereignisbaums — Utiliser l'arbre d'événements**2. Beispiel: • Exemple 2:**

Urne mit 6 Kugeln: 2 rote (r), 4 blaue (b). Ziehe 2 mal eine Kugel ohne zurücklegen.

• *Urne qui contient 6 boules: 2 rouges r), 4 bleues (b). Tirer 2 boules sans les remettre.*

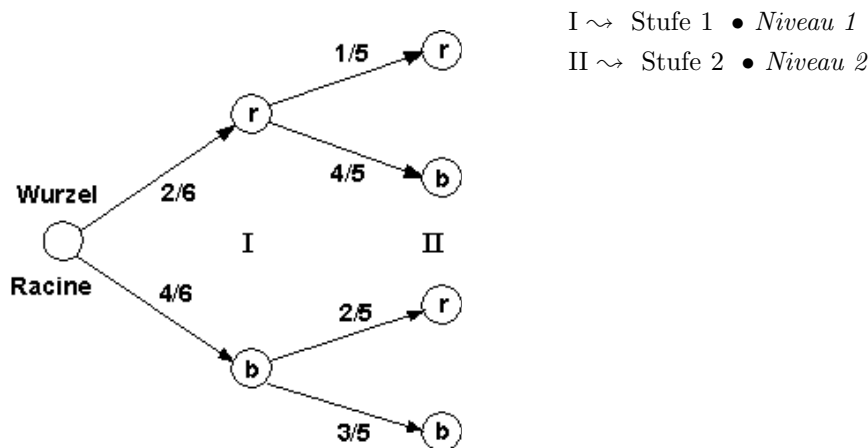
1. A: Wahrscheinlichkeit, 2 gleichfarbige Kugeln zu ziehen?

• *Probabilité qu'on tire 2 fois une boule avec la même couleur?*

2. B: Wahrscheinlichkeit, 2 verschiedenfarbige Kugeln zu ziehen?

• *Probabilité qu'on tire 2 fois une boule avec une couleur différente?*

Lösung: • Solution: Benutze einen Ereignisbaum: • *Utiliser un arbre d'événements:*

**Wichtig: • Important: Methode • Méthode**

Beachte: Den Pfeilen (Pfaden) entlang handelt es sich um abhängige Ereignisse, d.h. um Bedingte Wahrscheinlichkeit. An einem Verzweigungspunkt wird eine Bedingung gesetzt für die folgende Stufe. Daher multiplizieren sich die Wahrscheinlichkeiten auf den Pfaden. Hingegen sind die Ereignisse zu einem gegebenen Niveau ausschliessend. Die Wahrscheinlichkeiten addieren sich.

• *Considérer: Le long des flèches (sentiers) il s'agit d'événements dépendants, c.-à.-d. de prbabilité conditionnelle. À un point de bifurcation, une condition est posée pour le niveau suivant. Par conséquent les probabilités se multiplient le long des sentiers. Par contre les événements à un niveau donné sont excluants. Les probabilités s'additionnent.*

Lösung: • **Solution:**

$$\leadsto P(A) = \frac{2}{6} \cdot \frac{1}{5} + \frac{4}{6} \cdot \frac{3}{5} = \frac{7}{15}, \quad P(B) = \frac{2}{6} \cdot \frac{4}{5} + \frac{4}{6} \cdot \frac{2}{5} = \frac{8}{15}, \quad P(A) + P(B) = 1$$

3. Beispiel: • **Exemple 3:**

Wie gross ist die Wahrscheinlichkeit, beim k -ten Wurf mit einer idealen Münze erstmals Zahl zu treffen? Was ist die Summe über k aller solcher Werte für die Wahrscheinlichkeit? • *Quelle est la probabilité de réaliser la première fois pile au k -ème essai en lançant une monnaie idéale? Quelle est la somme sur k de toutes ces valeurs de probabilité?*

$$P(A_k) = \left(\frac{1}{2}\right)^k, \quad \sum_{k=1}^{\infty} P(A_k) = \sum_{k=1}^{\infty} \left(\frac{1}{2}\right)^k = \left(\sum_{k=0}^{\infty} \left(\frac{1}{2}\right)^k\right) - 1 = \frac{1}{1 - \frac{1}{2}} - 1 = 2 - 1 = 1$$

4.5.4 Totale Wahrscheinlichkeit — Probabilité totale

Bsp.: • **Exemple:**

Bei einem Zufallsexperiment ziehen wir aus einer von drei gegebenen Urnen U_i eine Kugel. U_1 enthält zwei rote und zwei blaue, U_2 zwei rote und eine blaue, U_3 eine rote und zwei blaue. • *Lors d'une expérience aléatoire, nous tirons d'une de trois urnes données U_i une boule. U_1 contient deux boules rouges et deux bleues, U_2 deux rouges et une bleue, U_3 une rouge et deux bleues.*

$\leadsto$ **Fragen:** • **Questions:**

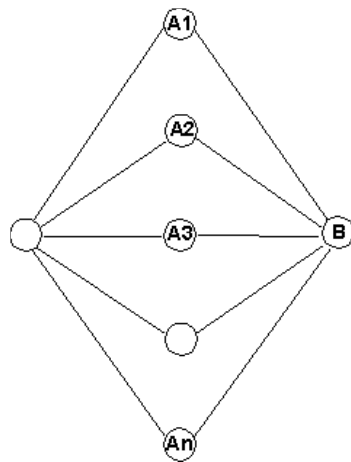
1. Wie gross ist die Wahrscheinlichkeit, eine rote Kugel zu ziehen?
• *Quelle est la probabilité de tirer une boule rouge?*
2. Wie gross ist die Wahrscheinlichkeit, dass die gezogene rote Kugel aus U_2 stammt?
• *Quelle est la probabilité que la boule rouge a été tirée de U_2 ?*

Lösung: • **Solution:**

1. Niveau 1: $P(\{U_i\}) = \frac{1}{3} \forall_i$
2. Niveau 2: $P(\{r\}|\{U_1\}) = \frac{2}{4}, \quad P(\{r\}|\{U_2\}) = \frac{2}{3}, \quad P(\{r\}|\{U_3\}) = \frac{1}{3}$
3. Niveau 1 $\rightarrow$ 2: $P(\{r \wedge U_1\}) = \frac{1}{3} \cdot \frac{2}{4} = \frac{1}{6}, \quad P(\{r \wedge U_2\}) = \frac{1}{3} \cdot \frac{2}{3} = \frac{2}{9}, \quad P(\{r \wedge U_3\}) = \frac{1}{3} \cdot \frac{1}{3} = \frac{1}{9}$
4. $\{r\}$: $P(\{r\}) = P(\{r \wedge U_1\}) + P(\{r \wedge U_2\}) + P(\{r \wedge U_3\}) = \sum_i P(\{U_i\}) \cdot P(\{r\}|\{U_i\})$
$$= \frac{1}{6} + \frac{2}{9} + \frac{1}{9} = \frac{1}{2} = 0.5$$

Selbstverständlich muss man bei diesem Experiment den Wert 0.5 erwarten, da total genau die Hälfte der Kugeln rot sind. Doch das Beispiel zeigt die Methode, die auch in andern Fällen funktioniert.

• *Évidemment on s'attend la valeur 0.5 lors de cette expérience, parce que en tout les boules rouges font exactement la moitié des boules. Mais l'exemple montre la méthode qui fonctionne aussi dans d'autres cas.*



Setzen wir B an Stelle von $\{r\}$ und A_i an Stelle von U_i , so erhalten wir die folgende Formel:

• Si nous mettons B à la place de $\{r\}$ et A_i à la place de U_i , nous obtenons la formule suivante:

Konsequenz: • **Conséquence:** Seien • Soient $i \neq k$

$$\Rightarrow A_i \cap A_k = \{\} \wedge A = (A_1 \cap B) \cup (A_2 \cap B) \cup \dots \cup (A_n \cap B), \quad P(B) = \sum_{i=1}^n P(A_i) \cdot P(B|A_i)$$

Definition: • **Définition:**

$P(B)$ heisst hier die **totale Wahrscheinlichkeit** für das Eintreten des Ereignisses B . • Ici $P(B)$ s'appelle la **probabilité totale** pour la réalisation de l'événement B .

Nun interessieren wir uns für die Wahrscheinlichkeit, dass der Pfad über A_k führt unter der Voraussetzung, dass B eingetreten ist. Die Wahrscheinlichkeit, dass der Pfad über A_k führt, ist dabei gleich der Wahrscheinlichkeit, dass A_k eintritt. Wir müssen demnach $P(A_k|B)$ berechnen. Gegeben sind aber $P(A_i)$ und $P(B|A_i)$.

• Maintenant nous nous intéressons à la probabilité que le sentier traverse A_k à condition que B soit réalisé. La probabilité que le sentier traverse A_k correspond à la probabilité que A_k se réalise. Nous devons donc calculer $P(A_k|B)$. $P(A_i)$ et $P(B|A_i)$ sont donnés.

Wir haben folgende Informationen zur Verfügung:

• Nous avons à la disposition les informations suivantes:

1. $P(B) = \sum_i P(A_i) \cdot P(B|A_i)$ (eben gewonnen) • (déduit en haut)
2. $P(A_k|B) = \frac{P(A_k \cap B)}{P(B)}$ (Multiplikationssatz) • (théorème de la multiplication)
3. $P(A_k \cap B) = P(A_k) \cdot P(B|A_k)$ (entlang einem Pfad) • (le long d'un sentier)

Konsequenz: • **Conséquence:** (Bayes'sche Formel) • **Formule de Bayes**

$$P(A_k|B) = \frac{P(A_k) \cdot P(B|A_k)}{\sum_i P(A_i) \cdot P(B|A_i)}$$

$$\leadsto P(\{U_2\}|\{r\}) = \frac{P(\{U_2\}) \cdot P(\{r\}|\{U_2\})}{\sum_{i=1}^3 P(\{U_i\}) \cdot P(\{r\}|\{U_i\})} = \frac{\frac{1}{3} \cdot \frac{2}{3}}{0.5} = \frac{2}{9 \cdot 0.5} = \frac{4}{9}$$

4.6 Wahrscheinlichkeitsverteilungen — Fonctions de répartition

4.6.1 Zufallsvariablen — Variables aléatoires

Wir betrachten Zufallsexperimente. Wenn wir aus einer Urne farbige Kugeln ziehen, so müssen wir der Farben erst eine Skala zuordnen, um sie auf der Achse eines Diagramms eintragen zu können und nachher mit Wahrscheinlichkeitsfunktionen als Zahlenfunktionen arbeiten zu können. Wenn wir hingegen mit einem Würfel würfeln, so sind die möglichen Resultate bereits durch Zahlen gegeben $\in \{1, 2, \dots, 6\}$.

• *Nous considérons des expériences de hasard. Si nous tirons des boules colorées d'une urne, nous devons d'abord attribuer une échelle graduée aux couleurs pour pouvoir les inscrire sur l'axe d'un diagramme à l'intention de travailler après avec les fonctions aléatoires comme fonctions sur les nombres. Si par contre nous jouons p. ex. aux dés avec un seul dé, les résultats possibles sont déjà les nombres $\in \{1, 2, \dots, 6\}$.*

Reelle Variablen können wir wie folgt als Funktionen interpretieren:

• *Nous pouvons interpréter les variable réelle comme fonctions comme il suit:*

$$x : z \mapsto x(z) = z.$$

Die Funktion x ordnet der Zahl z den Wert $z = w(z)$ zu, was selbstverständlich immer zutrifft und daher nie so erwähnt wird.

• *La fonction x applique le nombre z à la valeur $z = w(z)$, ce qui est toujours évident et par conséquent jamais mentionné.*

In analoger Weise wird nach der Absicht im Experiment durch eine **Zufallsvariable** einem Ergebnis $\omega \in \Omega$ (oder einem Elementarereignis $\{\omega\} \in \mathcal{A}$) ein Wert $X(\omega) \in \mathbb{R}$ zugeordnet.

• *De façon analogue on applique, d'après l'intention dans l'expérience, une valeur $X(\omega) \in \mathbb{R}$ par une variable aléatoire à un événement $\omega \in \Omega$ (ou un événement élémentaire $\{\omega\} \in \mathcal{A}$).*

1. Beispiel: • Exemple 1:

In einem Zufallsexperiment wird mit zwei Würfeln gewürfelt. Die Ergebnismenge ist $\Omega = \{(m, n) \mid m, n \in \{1, 2, \dots, 6\}\}$. Die Zufallsvariable X sei nun wie folgt definiert: $X((m, n)) = m + n \in \mathbb{R}$. Es gilt dann: $X((m, n)) \in \{2, 3, \dots, 12\}$. Die Menge der Ereignisse mit $X((m, n)) < 4$ ist dann $\{(1, 1), (1, 2), (2, 1)\}$.

• *Dans une expérience de hasard, on joue aux dés avec deux dés. L'ensemble des résultats est $\Omega = \{(m, n) \mid m, n \in \{1, 2, \dots, 6\}\}$. Maintenant la variable aléatoire X soit définie comme suit: $X((m, n)) = m + n \in \mathbb{R}$. Par conséquent il vaut: $X((m, n)) \in \{2, 3, \dots, 12\}$. L'ensemble des événements avec $X((m, n)) < 4$ est donc $\{(1, 1), (1, 2), (2, 1)\}$.*

2. Beispiel: • Exemple 2:

In einem Zufallsexperiment wird mit einem Würfel gewürfelt. Die Ergebnismenge ist $\Omega = \{1, 2, \dots, 6\}$. Die Zufallsvariable X sei nun wie folgt definiert: $X(n) = n$. Es gilt dann: $X(\omega_k) = \omega_k \in \mathbb{R}$.

• *Dans une expérience de hasard, on joue aux dés avec un dé. L'ensemble des résultats est $\Omega = \{1, 2, \dots, 6\}$. Maintenant la variable aléatoire X soit définie comme suit: $X(n) = n$. Par conséquent il vaut: $X(\omega_k) = \omega_k \in \mathbb{R}$.*

3. Beispiel: • Exemple 3:

In einem Zufallsexperiment werden Damen eines gegebenen Kollektivs nach dem Körperumfang ihres Traumpartners befragt. Für die Zufallsvariable gilt dann:

$X : \text{Dame} \mapsto X(\text{Dame}) = \text{Umfang Traumpartner} \in \mathbb{R}$

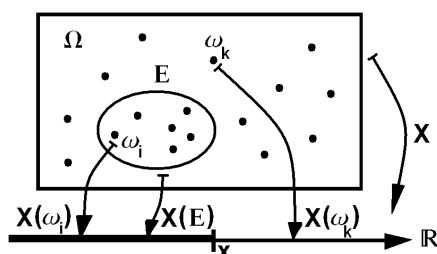
• *Dans une expérience de hasard, on questionne des dames d'un collectif donné sur le périmètre du corps de leur partenaire idéal. Pour la variable aléatoire il vaut donc:*

$$X : \text{Dame} \mapsto X(\text{Dame}) = \text{Umfang Traumpartner} \bullet \text{périmètre partenaire idéal} \in \mathbb{R}$$

Bemerkung: • **Remarque:**

Im 1. Beispiel ist $\{X(\dots)\} \subset \mathbb{R}$ diskret oder endlich, im 3. Beispiel ist $\{X(\dots)\} \subset \mathbb{R}$ theoretisch unendlich.

• Dans le 1er exemple $\{X(\dots)\} \subset \mathbb{R}$ est discret ou fini, dans le 3ème exemple $\{X(\dots)\} \subset \mathbb{R}$ est théoriquement infini.



Für unseren weiteren Gebrauch definieren wir nun die Zufallsvariable etwas genauer. Gegeben sei ein Zufallsexperiment mit einer Fundamentalmenge Ω und einer Ereignisalgebra $\mathcal{A}$, die eine σ -Algebra ist.

• Pour notre emploi à venir, nous définissons maintenant la variable aléatoire plus exactement. Soit donnée une expérience aléatoire avec un ensemble fondamental Ω mené d'une algèbre d'événement $\mathcal{A}$ qui est une σ -algèbre.

Definition: • **Définition:**

Eine Funktion $X : \Omega \mapsto \mathbb{R}$ mit $\omega \mapsto X(\omega)$ heisst **Zufallsvariable** oder **stochastische Variable**, **stochastische Funktion**, wenn gilt:

• Une fonction $X : \Omega \mapsto \mathbb{R}$ avec $\omega \mapsto X(\omega)$ s'appelle **variable aléatoire** ou **variable stochastique**, **fonction stochastique**, s'il vaut:

1. $\forall_{x \in \mathbb{R}} \{\omega \in \Omega \mid X(\omega) \leq x\} \in \mathcal{A}$
2. Zu $\mathcal{A}$ ist eine Wahrscheinlichkeitsfunktion P definiert, die den Wahrscheinlichkeitsaxiomen genügt. • Pour $\mathcal{A}$ soit définie une fonction de probabilité P , qui suffit aux axiomes de la probabilité.

Die erste Bedingung ist Messbarkeitsbedingung. Sie bedeutet, dass die Menge aller Ergebnisse, deren Realisation unterhalb eines bestimmten Wertes liegt, ein Ereignis bilden muss.

• La première condition est la condition de la mesurabilité. Elle signifie que l'ensemble de tous les résultats dont la réalisation possède une valeur au-dessous d'une valeur x donnée, est considéré comme événement.

Bemerkung: • **Remarque:**

Da es sich bei $\{\omega\}$ und somit auch bei $\{X(\omega)\}$ um Zufallsergebnisse handelt, sagen wir: Der Wert von X **hängt vom Zufall ab**.

• Comme à $\{\omega\}$ et donc aussi à $\{X(\omega)\}$ il s'agit d'événements de hasard, nous disons: La valeur de X **dépend du hasard**.

Schreibweise: • **Façon d'écrire:** $X \leq x := \{\omega \in \Omega \mid X(\omega) \leq x\}$

$\leadsto X \leq x, x < X$ u.s.w. sind Mengen! • $X \leq x, x < X$ etc. sont des ensembles!

Sprechweise: • **Façon de dire:** Z.B. für $a \leq X \leq b$ sagen wir: X nimmt einen Wert aus $[a, b]$ an.

• P. ex. pour $a \leq X \leq b$ nous disons: X prend une valeur dans $a \leq X \leq b$.

Da Teilmengen von Ω Ereignisse sind, so folgt auf Grund der Definition von X und der Auswahl der Elemente sofort der Satz:

• Comme les sous-ensembles de Ω sont des événements, on déduit s'appuyant sur la définition de X et la manière de choisir les éléments tout de suite le théorème:

Satz: • **Théorème:**

1. $-\infty < X < \infty$ ist das sichere Ereignis. • *est l'événement sûr.*
2. Folgende Mengen sind immer Ereignisse: • *Les ensembles suivants sont toujours des événements: $a < X < b$, $a \leq X < b$, $a < X \leq b$, $a \leq X \leq b$, $a < X$, $a \leq X$, $X < a$, $x \leq a$*
3. $P(X \leq a) + P(a < X) = P(-\infty < X < \infty) = 1$

Konsequenz: • **Conséquence:** $P(a < X) = 1 - P(X \leq a)$

Bemerkung: • **Remarque:**

Wie wir oben in den Beispielen gesehen haben, müssen wir unterscheiden zwischen **diskreten** und **kontinuierlichen Zufallsvariablen**. Es ist zu erwarten dass bei uns im Falle der kontinuierlichen Zufallsvariablen die Wertemengen Intervalle sind. Aus Gründen der Mächtigkeit von Intervallen und der Messgenauigkeit macht in diesem Falle eine Aussage wie z.B. $P(X = a)$ praktisch nicht viel Sinn, im Gegensatz zu einer Aussage wie $P(X < a)$.

• *Comme nous avons vu dans les exemples précédents, nous devons distinguer entre variables aléatoires discrètes et continues. On s'attend à ce que chez nous dans le cas des variables aléatoires continues les domaines de valeurs sont des intervalles. Pour les raisons de puissance d'intervalle et de l'exactitude des measurements, une proposition comme p. ex. $P(X = a)$ ne fait pratiquement pas beaucoup de sens, par contre une proposition comme $P(X < a)$*

3. Beispiel: • **Exemple 3:** Würfeln mit einem Würfel, X = Zahl:

• *Jouer aux dés avec un dé, X = pile (nombre):*

$$P(X = 1) = P(X = 2) = \dots = P(X = 6) = \frac{1}{6}$$

$$\Rightarrow P(1 < X < 2) = 0, P(1 \leq X < 6) = \frac{5}{6}, P\left(\frac{3}{2} < X < \frac{5}{2}\right) = \frac{1}{6}$$

4. Beispiel: • **Exemple 4:** Urne, 4 rote und 6 blaue Kugeln. 2 ziehen ohne zurücklegen. X (gezogene Kugeln) = Anzahl gezogene rote Kugeln. • *Urne, 4 boules rouges et 6 boules bleues. En tirer 2, sans les remettre. X (boules tirées) = nombre de boules rouges tirées.*

$$P(X = 0) = \frac{6}{10} \cdot \frac{5}{9} = \frac{1}{3}$$

$$P(X = 1) = ? \quad r_1 \rightsquigarrow P(Z_1) := P(X = 1|r_1) = \frac{4}{10} \cdot \frac{6}{9} = \frac{4}{15}$$

$$r_2 \rightsquigarrow P(Z_2) := P(X = 1|r_2) = \frac{6}{10} \cdot \frac{4}{9} = \frac{4}{15}$$

$$Z_1 \cap Z_2 = \{\} \Rightarrow P(Z_1 \cup Z_2) = \frac{4}{15} + \frac{4}{15} = \frac{8}{15}$$

$$P(X = 2) = \frac{4}{10} \cdot \frac{3}{9} = \frac{2}{15}$$

$$\begin{aligned}
&\leadsto P((X=0) \vee (X=1) \vee (X=2)) = \frac{1}{3} + \frac{8}{15} + \frac{2}{15} = 1 \\
&P(1 < X < 2) = 0 \\
&P(X \leq 1) = \frac{1}{3} + \frac{8}{15} = \frac{13}{15} \\
&P(X \geq 1) = \frac{8}{15} + \frac{2}{15} = \frac{2}{3} \\
&P(0.5 < X < 10) = P(X \geq 1) = \frac{2}{3}
\end{aligned}$$

4.6.2 Verteilungsfunktion — Fonction de répartition

Anlässlich der Erklärung der Zufallsvariable ist auf den Begriff der Wahrscheinlichkeitsfunktion verwiesen worden: Zu $\mathcal{A}$ ist eine Wahrscheinlichkeitsfunktion P definiert, die den Wahrscheinlichkeitsaxiomen genügt. Über das methodische Vorgehen zur Gewinnung von P ist aber nichts Allgemeines gesagt worden. Daher wollen wir hier zuerst den Begriff der Wahrscheinlichkeitsverteilung präzisieren:

• *Lors de l'explication de la variable aléatoire, on a mentionné l'idée de la fonction de probabilité: Pour $\mathcal{A}$ soit définie une fonction de probabilité P , qui suffit aux axiomes de la probabilité. Mais la méthode d'obtenir P de façon précise n'a pas été traitée. C'est pourquoi nous voulons d'abord préciser l'idée de la fonction de répartition:*

Gegeben sei ein Zufallsexperiment und dazu eine Zufallsvariable X . P sei nun praktisch so gegeben, dass zu jedem $x \in \mathbb{R}$ die Wahrscheinlichkeit $P(X \leq x)$ exakt definiert ist. Wichtig ist dabei die Einsicht, dass $P(X \leq x)$ für diskrete und für kontinuierliche Variablen problemlos vernünftig definierbar ist, während z.B. $P(X = x)$ im kontinuierlichen Fall wenig praktischen Sinn macht.

• *Soit donnée une expérience de hasard et liée à cela une variable de hasard X . Soit maintenant donné P de façon concrète et pratique qu'à chaque $x \in \mathbb{R}$ la probabilité $P(X \leq x)$ soit exactement définie. C'est important de comprendre que $P(X \leq x)$ peut être défini sans problème et raisonnablement pour des variables discrètes et continues tandis que par exemple dans le cas continu $P(X = x)$ ne donne pas de sens pratique.*

Definition: • **Définition:**

$F(x) = P(X \leq x)$ heisst **Verteilungsfunktion** von X .

• *s'appelle **fonction de répartition** de X .*

Bemerkung: • **Remarque:**

1. Nach Konstruktion ist $D_F = \mathbb{R}$.
• *D'après la construction il est $D_F = \mathbb{R}$.*
2. Mit Hilfe der Verteilungsfunktion können wir die „Wahrscheinlichkeitsfunktion“ oder „Wahrscheinlichkeitsverteilung“ im diskreten und im kontinuierlichen Fall auf verschiedene Arten definieren.
• *A l'aide de la fonction de répartition, nous pouvons définir la "fonction de probabilité" dans le cas discret et dans le cas continu de manière différente.*

Nach Definition ist die Verteilungsfunktion eine Wahrscheinlichkeit. Daher kann man ablesen:

• *D'après la définition, la fonction de répartition est une probabilité. C'est pourquoi on peut déduire:*

1. $\forall x : 0 \leq F(x) \leq 1$
2. $F(\infty) := \lim_{x \rightarrow \infty} F(x) = 1$

3. $F(-\infty) := \lim_{x \rightarrow -\infty} F(x) = 0$
4. $x_1 < x_2 \Rightarrow F(x_1) \leq F(x_2)$
5. $P(x_1 < X \leq x_2) = F(x_2) - F(x_1)$

4.6.3 Verteilungstypen — Types de répartitions

Problem: • **Problème:** Die Ausgangsfrage ist: Wie verteilen sich bei Zufallsexperimenten die verschiedenen Wahrscheinlichkeiten auf die verschiedenen Ereignisse? Vom Resultat her kennen wir schon diskrete und kontinuierliche Zufallsvariablen. Man kann aber auch die zugehörigen Experimente unterscheiden. Abgesehen von den Mischfällen unterscheiden wir zwei Klassen:

• *On est parti de la question: Comment est-ce que, lors d'expériences de hasard, les probabilités différentes sont distribuées aux événements différents? Partant du résultat, nous connaissons déjà des variables aléatoires discrètes et continues. Mais on peut aussi distinguer les expériences affiliés. Hormis les cas mixtes, nous distinguons deux classes:*

Experimentklassen: • *Classes d'expériences:*

1. **Zählexperimente** $\leadsto$ diskrete Verteilung • **Expériences de conter** $\leadsto$ *répartition discrète*
2. **Messexperimente** $\leadsto$ kontinuierliche Verteilung • **Expériences de mesurer** $\leadsto$ *répartition continue*

1. Beispiel: • Exemple 1: Volkszählung: • Recensement de la population:

Verteilung der Bewohner auf die Gemeinden $\leadsto H(X = a) = \text{Anzahl Bewohner der Gemeinde } a$.

Hier wird die Verteilung durch eine Wahrscheinlichkeitsfunktion $\frac{H(X = a)}{H(S)} = P(X = a)$ beschrieben.

• *Répartition des habitants sur les communes* $\leadsto H(X = a) = \text{nombre d'habitants de la commune } a$.

Ici, la répartition est décrite par une fonction de probabilité $\frac{H(X = a)}{H(S)} = P(X = a)$.

2. Beispiel: • Exemple 2: Gewichtsfunktion: • Fonction de poids:

Verteilung des Gewichts auf Balken gleichen Querschnitts $\leadsto H(X \leq x) = \int_a^x f(t) dt = \text{Anzahl Balken der Länge } x \in [a, b]$.

Hier wird die Verteilung durch eine Wahrscheinlichkeitsdichtefunktion $f(x)$ beschrieben:

$\leadsto f(x) = \frac{d(\int_a^x f(t) dt)}{dx}$. f liefert ein Modell für die Grundgesamtheit, die oft gar nicht genau bekannt ist.

• *Répartition du poids sur des poutres à la coupe transversale égale* $\leadsto H(X \leq x) = \int_a^x f(t) dt = \text{nombre de poutres de la longueur } x \in [a, b]$.

Ici, la répartition est décrite par une fonction de densité de probabilité $f(x)$:

$\leadsto f(x) = \frac{d(\int_a^x f(t) dt)}{dx}$. f livre un modèle pour l'ensemble de base, qui souvent n'est pas exactement connu.

Zwischen Grundgesamtheit und Modell können wir folgenden Vergleich anstellen:

• *Nous pouvons comparer l'ensemble de base et le modèle comme il suit:*

Grundgesamtheit • *Ensemble de base*

1. Wirklichkeit • *Réalité*
2. Praxis: Häufigkeitsverteilung oder Dichte • *Pratique: Répartition de fréquence ou densité*
3. Oft nur Stichproben • *Souvent seulement échantillons*

Modell • *Modèle*

1. Theorie • *Théorie*
2. Wahrscheinlichkeitsverteilung • *Répartition de probabilité*
3. Problem der Bonität des Modells (Tests) • *Problème de qualité du modèle (tests)*

4.6.4 Diskrete Verteilung — Répartition discrète**Wahrscheinlichkeitsfunktion — Fonction aléatoire**

Für diskrete Variablen sind endliche Definitionsbereiche nicht zwingend. Wir definieren daher:

• *Pour les variables discrètes, les domaines de définition ne doivent pas nécessairement être finis. Nous définissons par conséquent:*

Definition: • **Définition:**

Eine Zufallsvariable X heisst **diskret**, wenn gilt:

• *Une variable aléatoire X s'appelle **discrète**, s'il vaut:*

1. X kann nur endlich oder abzählbar unendlich viele Werte x_i mit $P(X = x_i) \neq 0$ annehmen.
• *X ne peut atteindre qu'un nombre fini ou infini de façon énumérable de valeurs x_i avec $P(X = x_i) \neq 0$.*
2. $\{x_i \mid P(X = x_i) \neq 0\}$ hat keine Häufungspunkte.
• *$\{x_i \mid P(X = x_i) \neq 0\}$ n'a pas de points d'accumulation (limite).*

Konsequenz: • **Conséquence:**

In jedem endlichen Intervall $\subset \mathbb{R}$ liegen nur endlich viele Werte x_i mit $P(X = x_i) \neq 0$.

• *Dans chaque intervalle fini $\subset \mathbb{R}$ il n'y a qu'un nombre fini de x_i avec $P(X = x_i) \neq 0$.*

Bezeichnungen: • **Notations:**

Statt $P(X = x_i)$ schreiben wir kurz $P(x_i)$. Die Werte x mit $P(x) \neq 0$ seien nummeriert: $x_1, x_2, x_3, \dots$. Die zugehörigen Wahrscheinlichkeiten schreiben wir ebenfalls kurz: $P(X = x_i) := P(x_i) := p_i$.

• *Au lieu d'écrire $P(X = x_i)$ nous notons brièvement $P(x_i)$. Les valeurs x avec $P(x) \neq 0$ soient numérotées: $x_1, x_2, x_3, \dots$. Pour les probabilités appartenantes à $x_1, x_2, x_3, \dots$ nous notons aussi brièvement: $P(X = x_i) := P(x_i) := p_i$.*

Definition: • **Définition:**

$$f(X) = \begin{cases} p_i & X = x_i \\ 0 & \text{sonst} \end{cases} \quad \bullet \text{ autrement}$$

heisst **Wahrscheinlichkeitsfunktion** der Zufallsvariablen X

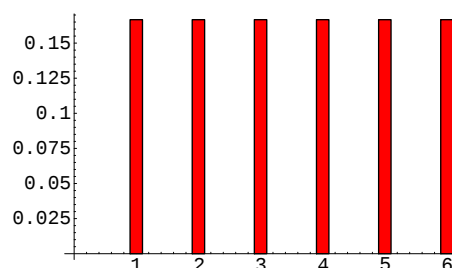
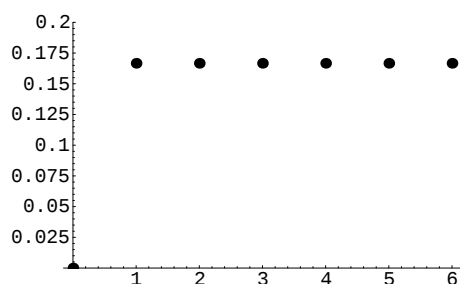
• *s'appelle **fonction de probabilité** de la variable aléatoire X*

Bemerkung: • **Remarque:**

Der Graph einer solchen Wahrscheinlichkeitsfunktion besteht aus isolierten Punkten. Die graphische Darstellung wird daher oft als unübersichtlich empfunden und daher durch ein **Stabdiagramm** ersetzt.

• *Le graphe d'une telle fonction de probabilité consiste en points isolés. Souvent cette représentation graphique donne une impression peu claire et par conséquent on la remplace par un diagramme de bâton.*

Bsp.: • **Exemple:** Würfel: • **Dé:**



Es ist unmittelbar einsichtig, dass durch die Wahrscheinlichkeitsfunktion die Wahrscheinlichkeitsverteilung oder die Verteilung der Zufallsvariablen X vollständig bestimmt ist. • *Il est bien clair que la fonction aléatoire détermine complètement la fonction de répartition ou la répartition de la variable aléatoire X .*

Satz: • **Théorème:**

1. $\sum_i f(x_i) = 1$
2. $P(a \leq X \leq b) = \sum_{a \leq x_i \leq b} f(x_i) = \sum_{a \leq x_i \leq b} p_i$

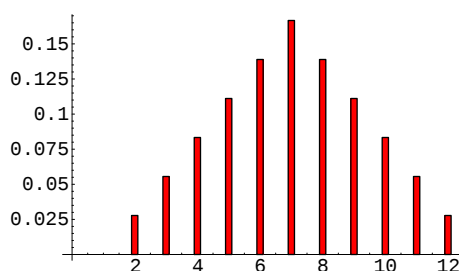
Bsp.: • **Exemple:** Bestimme die Wahrscheinlichkeitsfunktion (Graph, Wertetabelle) beim Zufallsexperiment „würfeln mit zwei Würfeln“.

• *Déterminer la fonction aléatoire (graphe, tableau des valeurs) pour l'expérience de hasard "jouer aux dés avec deux dés".*

Wertetabelle: • *Tableau des valeurs:*

x	2	3	4	5	6	7	8	9	10	11	12
$f(x)$	$\frac{1}{36}$	$\frac{2}{36}$	$\frac{3}{36}$	$\frac{4}{36}$	$\frac{5}{36}$	$\frac{6}{36}$	$\frac{5}{36}$	$\frac{4}{36}$	$\frac{3}{36}$	$\frac{2}{36}$	$\frac{1}{36}$

Graph: • *Graphe:*



Verteilungsfunktion — Fonction de répartition

Sei X eine Zufallsvariable aus einem Zufallsexperiment. Sei $P(X \leq x)$ damit bestimmt, $x \in \mathbb{R}$ beliebig. Sei $I = \{X \mid X \leq x\}$. Damit schreiben wir kurz: $P(X \leq x) = P(I)$. Von Seite 84 wissen wir, dass damit durch $F(x) := P(X \leq x)$ die Verteilungsfunktion der Zufallsvariablen X für alle $x \in \mathbb{R}$ definiert ist.

• Soit X une variable aléatoire d'une expérience de hasard. Soit $P(X \leq x)$ donc défini, $x \in \mathbb{R}$ quelconque. Soit $I = \{X \mid X \leq x\}$. Ainsi nous écrivons brièvement: $P(X \leq x) = P(I)$. De la page 84 on sait qu'ainsi la fonction aléatoire de la variable aléatoire X est définie par $F(x) := P(X \leq x)$ pour tous les $x \in \mathbb{R}$.

Bsp.: • **Exemple:** Wurf einer Münze, einmal. $X = X(\{\text{Anzahl Kopf}\})$. • Jeter une fois une monnaie. $X = X(\{\text{nombre de faces (têtes) réalisées}\})$.

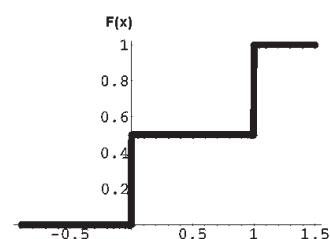
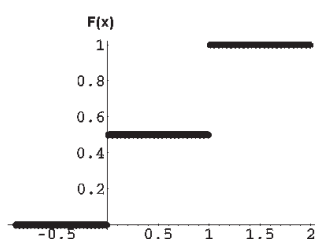
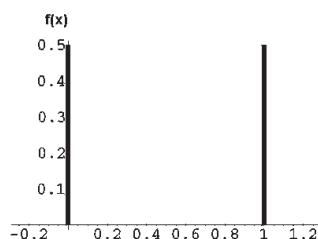
Es ist: • On trouve: $P(X = 0) = P(X = 1) = 0.5$.

1. $x < 0 \leadsto P(X < 0) = 0 \Rightarrow F(x) = 0, x < 0$
2. $x = 0 \leadsto P(X = 0) = 0.5 \Rightarrow F(0) = P(x \leq 0) = 0.5$
3. $x \in (0, 1) \leadsto \dots \Rightarrow F(x) = 0.5, x \in (0, 1)$
4. $x = 1 \leadsto P(X \leq 1) = P(X < 1) + P(X = 1) = P(X = 0) + P(X = 1) = 0.5 + 0.5 = 1$
5. $x > 1 \leadsto \dots \Rightarrow P(x) = 1, x > 1$

Resultat: • Résultat:

x	$0 < x$	$0 \leq x < 1$	$1 \leq x$
$F(x)$	0	0.5	1

Diagramme: • Diagrammes:



Problem: • **Problème:** **Geg.:** • **Donné:** $X \in (a, b] \leadsto a < X \leq b$

Frage: • **Question:**

Wie lässt sich aus $F(x)$ die Wahrscheinlichkeit $P(a < X \leq b)$ berechnen?

• Comment est-ce qu'on peut calculer la probabilité $P(a < X \leq b)$ de $F(x)$?

Lösung: • **Solution:**

$$\{X \mid X \leq a\} \cup \{X \mid a < X \leq b\} = \{X \mid X \leq b\} \Rightarrow P(X \leq a) + P(a < X \leq b) = P(X \leq b)$$

$$P(X \leq a) = F(a), P(X \leq b) = F(b) \Rightarrow \{X \mid X \leq b\} \Rightarrow F(a) + P(a < X \leq b) = F(b)$$

Damit lässt sich bei der diskreten Verteilung die Wahrscheinlichkeitsfunktion aus der Verteilungsfunktion berechnen. Den folgenden Satz kennen wir schon von früher (Seite 84):

• *Donc si la répartition est discrète, on peut calculer la fonction aléatoire de la fonction de répartition. Nous connaissons le théorème suivant depuis en haut (page 84):*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Diskrete Verteilung
• *Répartition discrète*

Beh.: • **Thè.:**

$$P(a < X \leq b) = F(b) - F(a)$$

Sei • *Soit* $x_1 \leq x_2 \leq \dots \leq x_n \leq x$

Von Seite 87 wissen wir: • *De la page 87 nous savons:*

$$P(a \leq X \leq b) = \sum_{a \leq x_i \leq b} f(x_i) = \sum_{a \leq x_i \leq b} p_i, \quad f(x_i) = p_i \rightsquigarrow F(x) = P(X \leq x) = \sum_{x_i \leq x} f(x_i) = \sum_{x_i \leq x} p_i \rightsquigarrow$$

Satz: • **Théorème:** $F(x) = \sum_{x_i \leq x} f(x_i) = \sum_{x_i \leq x} p_i$

Bsp.: • **Exemple:** $X = X(\{\text{Augenzahl beim Würfeln mit einem Würfel.}\})$

• $X = X(\{\text{Nombre de points en jouant avec un dé.}\})$

$F(x)$ ist weniger anschaulich als $f(x)$, besitzt jedoch Vorteile in der Handhabung.

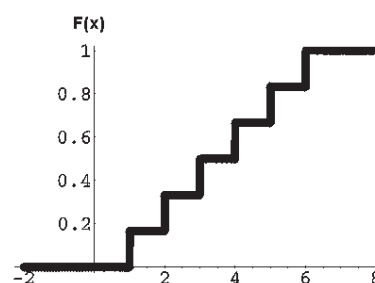
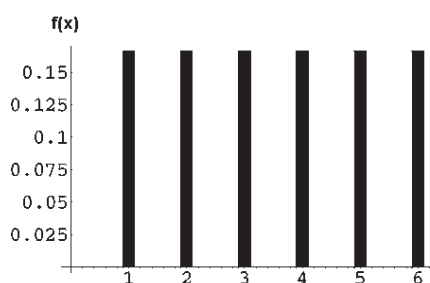
• $F(x)$ est moins compréhensible que $f(x)$, mais F possède des avantages pour l'application.

Es ist: • *On trouve:* $P(X = 1) = P(X = 2) = \dots = P(X = 6) = \frac{1}{6}$.

$\rightsquigarrow$ **Resultat:** • **Résultat:**

x	$x < 1$	$1 \leq x < 2$	$2 \leq x < 3$	$3 \leq x < 4$	$4 \leq x < 5$	$5 \leq x < 6$	$6 \leq x$
$F(x)$	0	$\frac{1}{6}$	$\frac{2}{6} = \frac{1}{3}$	$\frac{3}{6} = \frac{1}{2}$	$\frac{4}{6} = \frac{2}{3}$	$\frac{5}{6}$	$\frac{6}{6} = 1$

Diagramme: • *Diagrammes:*



$\rightsquigarrow$ Z.B. • *P.ex.* $P(X = 2 \vee X = 3) = P(1 < X \leq 3) = F(3) - F(1) = \frac{1}{2} - \frac{1}{6} = \frac{1}{3}$

Bemerkung: • **Remarque:**

Eine Funktion wie $F(x)$ im obigen Beispiel nennen wir **Treppenfunktion**. Sie besitzt an diskreten Stellen (mögliche Werte x_i der Zufallsvariablen X) Sprünge, dazwischen verläuft sie konstant. Die Sprünge entstehen, weil dort $p(x) \neq 0$ ist. Treppenfunktionen kann man mit Hilfe der „Sprungfunktion“ $h(x)$ (vgl. Dirac'sche δ -Funktion, vgl. Analysis) konstruieren:

• *Une fonction comme $F(x)$ dans l'exemple susdit s'appelle **fonction en escalier (à palier)**. Elle a des sauts à places discrètes (valeurs x_i possibles de la variable aléatoire X), entre ces places elle est constante. Les sauts résultent de la valeur $p(x) \neq 0$. On peut construire la fonction en escalier à l'aide de la "fonction de saut" $h(x)$ (voir fonction δ de Dirac, voir analyse):*

Satz: • **Théorème:**

$$F(x) = \sum_{k=1}^n h(x - x_k) \cdot p(x_k), \quad p(x_k) \neq 0, \quad k = 1, \dots, n, \quad h(x) = \operatorname{sgn}\left(\frac{\operatorname{sgn}(x) + 1}{2}\right)$$

4.6.5 Kontinuierliche Verteilung — Répartition continue

Das Problem der Wahrscheinlichkeitsfunktion — Le problème de la fonction aléatoire

Wir definieren für unsere Bedürfnisse: • *Pour nos besoins nous définissons:*

Definition: • **Définition:**

Eine Zufallsvariable X heisst **stetig in einem Intervall I** , wenn alle Werte in I angenommen werden können.

• *Une variable aléatoire X s'appelle **continue dans un intervalle I** , si toutes les valeurs de I peuvent être réalisées.*

Problem: • **Problème:** Macht bei einer stetigen Zufallsvariablen $P(X = x_i)$ noch einen Sinn?

• *Est-ce que pour une variable aléatoire continue $P(X = x_i)$ a un sens?*

Das Problem ist ja, dass in diesem Fall unendlich viele Werte $x_i \in I$ angenommen werden können. Da aber im diskreten Fall die aufsummierte Wahrscheinlichkeit 1 ist, müsste daher auch hier konsequenterweise für die meisten x_i dann $P(X = x_i) = 0$ sein. Das besagt aber, dass diese angenommenen Werte dann doch nicht angenommen werden. Wir landen also in einem Widerspruch. Daher ist es in unserem Rahmen nicht sinnvoll, nach der Wahrscheinlichkeit $P(X = x_i)$ und daher nach einer Wahrscheinlichkeitsfunktion zu fragen. Wir fragen besser nach der Wahrscheinlichkeit für Intervalle, z.B. $P(a \leq X < b) = ?$ Das ist bei einer gegebenen Verteilungsfunktion kein Problem.

• *Le problème est que dans ce cas il y a un nombre infini de valeurs $x_i \in I$ qui peuvent être réalisées. Comme dans le cas discret la somme de la probabilité est 1, logiquement et par conséquent aussi dans ce cas pour la plupart des x_i il vaut $P(X = x_i) = 0$. Ça signifie que ces valeurs réalisées par contre ne sont pas réalisées. Ainsi nous arrivons à une contradiction. Par conséquent, dans ce cadre il n'est pas raisonnable, de demander la probabilité $P(X = x_i)$ ni par conséquent une fonction de probabilité. Nous demandons donc plutôt la probabilité pour les intervalles, par exemple $P(a \leq X < b) = ?$ Ceci n'est pas un problème si une fonction de probabilité est donnée.*

Situation beim diskreten Fall — Situation au cas discret

Sei also: • *Soit donc:* $F(x) = P(X \leq x)$

Im diskreten Fall ist: • *Dans le cas discret il est:*

$$F(x) = P(X \leq x) = \sum_{x_i \leq x} f(x_i), \quad x_i : x_1 < x_2 < x_3 < \dots$$

Da die diskreten Werte x_i geordnet sind, besteht eine bijektive Beziehung zwischen dem Wert x_i und dem Index i :

• *Comme les valeurs discrètes x_i sont ordonnées, il existe une liaison bijective entre la valeur x_i et l'index i :*

$$\varphi : x_i \mapsto \varphi(x_i) = i$$

Damit existiert auch $\varphi^{-1} = \psi$: • *Il existe donc aussi $\varphi^{-1} = \psi$:*

$$\psi(i) = x_i$$

$$\begin{aligned} \text{Sei } & \bullet \text{ Soit } \varphi(x_j) = j, \quad \psi(j) = x_j \Rightarrow F(x_j) = \sum_{x_i \leq x_j} f(x_i) = \sum_{\psi(i) \leq \psi(j)} f(\psi(i)) := \\ & = \sum_{\psi(i) \leq \psi(j)} \tilde{f}(i) = \sum_{\psi(i) \leq \psi(j)} \tilde{f}(i) \cdot 1 = \sum_{\psi(i) \leq \psi(j)} \tilde{f}(i) \cdot \underbrace{(i+1-i)}_{\Delta i} = \sum_{\psi(i) \leq \psi(j)} \tilde{f}(i) \cdot \Delta i \end{aligned}$$

$\tilde{f}(i) \cdot \Delta i$ hat die Bedeutung eines Rechteckinhalts. $F(x)$ hat somit die Bedeutung eines Flächeninhalts unter einer Treppenkurve. Die Bedeutung des Flächeninhalts hat aber auch im stetigen Fall noch Gültigkeit.

• $\tilde{f}(i) \cdot \Delta i$ a la signification d'une aire d'un rectangle. $F(x)$ a donc la signification d'une aire sous une courbe en escalier. Cette signification d'une aire garde sa valabilité aussi pour le cas continu.

Wahrscheinlichkeitsdichte — Densité de probabilité

Im stetigen Fall liegt es daher nahe, $F(x) = P(X \leq x)$ ebenfalls als Flächeninhalt unter einer Kurve zu definieren. • *Il est donc bien clair que dans le cas continu, on peut définir $F(x) = P(X \leq x)$ comme mesure de surface sous une courbe.*

$$F(x) = P(X \leq x) = \int_{-\infty}^x f(x^*) dx^* \Rightarrow f(x) = \frac{dF}{dx}$$

$f(x) = \frac{dF}{dx}$ kann hier nicht mehr die Bedeutung einer Wahrscheinlichkeit haben, denn alle Wahrscheinlichkeiten müssen sich ja zur Gesamtwahrscheinlichkeit 1 aufsummieren lassen können. Daher redet man hier von einer **Wahrscheinlichkeitsdichte**.

• *Ici, $f(x) = \frac{dF}{dx}$ ne peut plus avoir la signification d'une probabilité, car toutes les probabilités doivent donner la somme 1. C'est pourquoi on parle ici d'une **densité de probabilité**.*

$$\leadsto P(a < x \leq b) = F(b) - F(a)$$

Diagramme: • *Diagrammes:*



Bsp.: • **Exemple:** Roulettepiel: • *Jouer à la roulette:*

$$0 \leq X < 2\pi$$

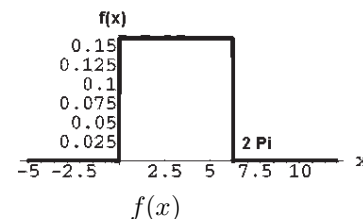
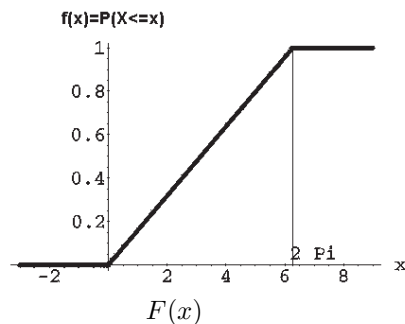
$X :=$ Winkel des Zeigers mit einer festen Richtung. • $X :=$ *angle entre l'aiguille et une direction fixe.*

Brauchbares Roulett • *Roulette utilisable* $\leadsto P(X \leq x) = \frac{x}{2\pi}$

$$\Rightarrow F(x) = \begin{cases} 0 & x \leq 0 \\ \frac{x}{2\pi} & 0 < x \leq 2\pi \\ 1 & 2\pi < x \end{cases}$$

$$f(x) = F'(x) \Rightarrow f(x) = \begin{cases} \frac{1}{2\pi} & x \in (0, 2\pi) \\ 0 & x \notin (0, 2\pi) \end{cases}$$

Diagramme: • *Diagrammes:*



Bemerkung: • **Remarque:**

Eine Verteilung wie im letzten Beispiel nennen wir **gleichförmige Verteilung** oder **Rechtecksverteilung**. • *Nous appelons une distribution comme on la voit dans le dernier exemple **distribution rectangulaire** ou **distribution uniforme**.*

Hinweis: • **Indication:**

Man beachte die Analogie zwischen Masseverteilung in der Physik und Wahrscheinlichkeitsverteilung. Punktmassen entsprechen dann Wahrscheinlichkeiten im diskreten Fall.

• *Il faut bien considérer l'analogie entre la distribution de la masse dans la physique et la distribution aléatoire. Les masses ponctuelles correspondent à la probabilité dans le cas discret.*

4.7 Mass- oder Kennzahlen einer Verteilung — Mesures de répartition

4.7.1 Allgemeines — Considérations générales

Problem: • **Problème:** Wie lassen sich zwei Verteilungen rasch vergleichen? • *Comment peut-on vite comparer deux distributions?*

1. Qualitativer Vergleich von $f(x)$, $F(x)$ oder Graphen (unendliche Mengen!).
 - *Comparaison qualitative de $f(x)$, $F(x)$ ou des graphiques (quantités infinies!).*

2. Quantitativer Vergleich charakteristischer Masszahlen: Mittelwert μ , Varianz σ^2 , Standardabweichung σ , Schiefe γ u.s.w.. • *Comparaison quantitative de mesures caractéristiques: La moyenne μ , variance σ^2 , écart type (quadratique moyen) (σ), dissymétrie γ etc..*

4.7.2 Mittelwert — Valeur moyenne

1. Diskreter Fall: • Cas discret:

Vergleich mit der Physik: Unter dem „Mittelpunkt einer Massenverteilung“ verstehen wir den Schwerpunkt. Betrachtet man an Stelle von Wahrscheinlichkeitsverteilungen nun statistische Häufigkeitsverteilungen, so tritt an die Stelle der Massenverteilung das arithmetische Mittel oder kurz der Mittelwert. Bei einer Wahrscheinlichkeitsverteilung tritt an die Stelle des Schwerpunktes somit auch Mittelwert. Es gilt:

• *Comparaison avec la physique: Par le "centre d'une distribution de masses", nous comprenons le centre de gravité. Si l'on considère des distributions de fréquence statistiques à la place de distributions aléatoires, la moyenne arithmétique (ou brièvement la moyenne) prend la place de la distribution de masse. Quant à une distribution aléatoire, c'est donc aussi la valeur moyenne qui remplace le centre de gravité. Il vaut:*

$$\bar{x} \cdot m = \sum_{i=1}^n x_i \cdot m_i \Rightarrow \bar{x} = \sum_{i=1}^n \frac{x_i \cdot m_i}{m} = \sum_{i=1}^n \frac{m_i}{m} \cdot x_i \hat{=} \sum_{i=1}^n p_i \cdot x_i := \sum_{i=1}^n f(x_i) \cdot x_i$$

Dabei kann $n \in \mathbb{N}$ oder $n = \infty$ sein. • *On peut trouver le cas $n \in \mathbb{N}$ ou $n = \infty$.*

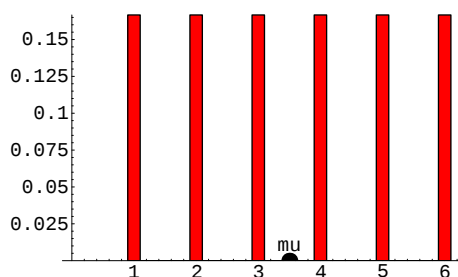
Daher definieren wir: • *C'est pourquoi nous définissons:*

Definition: • Définition: Mittelwert: • **Valeur moyenne:** $\mu := \sum_i f(x_i) \cdot x_i \cdot 1$

Konsequenz: • Conséquence:

$$(a) \ i \in \{1, 2, \dots, n\}, \ f(x_i) = p_i \Rightarrow \mu = \sum_{i=1}^n p_i \cdot x_i$$

$$(b) \ p_i = \frac{1}{n} \Rightarrow \mu = \frac{1}{n} \sum_{i=1}^n x_i$$



Bsp.: • Exemple:

X = Augenzahl beim Würfeln.

• X = nombre de points en jouant avec un dé.

$$f(x_i) = \frac{1}{6} \quad x_i = 1, 2, 3, \dots, 6 \Rightarrow 1 \cdot \frac{1}{6} + 2 \cdot \frac{1}{6} + \dots + 6 \cdot \frac{1}{6} = \frac{1}{6} \cdot \sum_{i=1}^6 i = \frac{1}{6} \cdot 21 = 3.5.$$

2. Übertragung des Begriffs auf den stetigen Fall: • Transcription de l'idée pour le cas discontinu:

Sei • Soit $x_i = i \Rightarrow \Delta x_i = (i+1) - i = 1 \Rightarrow$

$$\mu = \sum_{i=n}^6 f(x_i) \cdot x_i = \sum_{i=n}^6 f(x_i) \cdot x_i \cdot 1 = \sum_{i=n}^6 f(x_i) \cdot x_i \cdot \Delta x_i \rightarrow \int_{\dots}^{\dots} f(x) \cdot x \cdot dx$$

Daher definieren wir analog zum Schwerpunkt der Physik:

• C'est pourquoi nous définissons analogiquement au centre de gravité de la physique:

Definition: • **Définition:** **Mittelwert:** • **Valeur moyenne:**

$$\mu := \int_{-\infty}^{\infty} f(x) \cdot x \cdot dx$$

4.7.3 Erwartungswert — Valeur d'espérance

Bsp.: • **Exemple:**

Zwei Spieler A und B spielen wie folgt:

Vor dem Spiel zahlt A an B den Betrag „durchschnittliche Gewinnerwartung beim Spiel“. A darf dann würfeln, während B an A Gewinne auszahlt nach dem folgenden Schema: Beim Resultat 1 oder 2 erhält A von B Fr. 10.-, beim Resultat 3 oder 4 erhält A von B Fr. 20.-, beim Resultat 5 erhält A von B Fr. 40.-, beim Resultat 6 erhält A von B Fr. 80.-.

Was ist die durchschnittliche Gewinnerwartung?

• Les deux joueurs A et B jouent comme il suit: Avant le jeu A paie à B le montant = profit moyen attendu au jeu. Alors, A peut jouer aux dés pendant que B paie le gain après le schéma suivant: Au résultat 1 ou 2, A reçoit de B sfr. 10. -, au résultat 3 ou 4, A reçoit de B sfr. 20. -, au résultat 5, A reçoit de B sfr. 40. -, au résultat 6, A reçoit de B sfr. 80. Quelle est le profit moyen?

$$\leadsto \bar{G} = 10 \cdot \frac{1}{6} + 10 \cdot \frac{1}{6} + 20 \cdot \frac{1}{6} + 20 \cdot \frac{1}{6} + 40 \cdot \frac{1}{6} + 80 \cdot \frac{1}{6} = 30 \quad \text{oder} \quad \bullet \quad \text{ou } 6 \cdot \bar{G} = 10 + 10 + 20 + 20 + 40 + 80$$

In diesem Beispiel sehen wir eine Zuordnung folgender Art:

• Dans cet exemple, nous constatons consécutivement un classement de la sorte suivante:

$X \longleftrightarrow g(X)$ resp. • resp. $x_i \longleftrightarrow g(x_i)$	X	1	2	3	4	5	6
	$g(X)$	10	10	20	20	40	80

X ist unabhängige Zufallsvariable, $g(X)$ ist abhängige Zufallsvariable.

• X est variable aléatoire indépendante, $g(X)$ est variable aléatoire dépendante.

Konsequenz: • **Conséquence:**

$g(X)$ ist ebenfalls Zufallsvariable • $g(X)$ est aussi variable aléatoire.

Um die durchschnittliche **Gewinnerwartung** zu berechnen, müssen wir daher x_i ersetzen durch $g(x_i)$, d.h. wir ersetzen $x_i \cdot f(x_i)$ durch $g(x_i) \cdot f(x_i)$. Dabei ist $f(x_i)$ die zu $g(x_i)$ gehörige Wahrscheinlichkeit.

• Pour calculer le gain moyen attendu (l'espérance mathématique), nous devons remplacer par conséquent x_i par $g(x_i)$, c.-à.-d. nous remplaçons $x_i \cdot f(x_i)$ par $g(x_i) \cdot f(x_i)$. Là, $f(x_i)$ est la probabilité appartenant à $g(x_i)$.

Wir definieren daher analog zur durchschnittlichen Gewinnerwartung ganz allgemein die **Erwartung** resp. den **Erwartungswert** $E(g(X))$ zur **Erwartungsfunktion** $g(X)$. Analog zur Definition zu μ setzen wir fest:

• *Nous définissons donc tout à fait analogiquement à l'espérance mathématique moyenne l'espérance (valeur d'espérance, valeur attendue) $E(g(X))$ pour la fonction d'espérance $g(X)$. Analogiquement à la définition de μ nous fixons:*

Definition: • **Définition:**

$$1. \text{ Diskreter Fall: } \bullet \text{ Cas discret: } E(g(X)) := \sum_i g(x_i) \cdot f(x_i) \cdot 1 = \sum_i g(x_i) \cdot f(x_i) \cdot \Delta x$$

$$2. \text{ Stetiger Fall: } \bullet \text{ Cas continu: } E(g(X)) := \int_{-\infty}^{\infty} g(x) \cdot f(x) dx$$

Bemerkung: • **Remarque:** Statt Erwartungswert sagen wir auch mathematischer Erwartungswert, Erwartung oder manchmal auch kurz Mittelwert.

• *Au lieu de dire valeur d'espérance, nous disons souvent aussi valeur d'espérance mathématique ou bien brièvement valeur moyenne.*

Konsequenz: • **Conséquence:** $g(x) = x \Rightarrow E(g(x)) = E(x) = \mu$

Speziell: • *Spécialement:* (Diskreter Fall) • (Cas discret)

$$1. E(g(X)) = \sum_{i=1}^n g(x_i) \cdot f(x_i) = \left\langle \begin{pmatrix} g(x_1) \\ \vdots \\ g(x_n) \end{pmatrix}, \begin{pmatrix} f(x_1) \\ \vdots \\ f(x_n) \end{pmatrix} \right\rangle = \langle \vec{g}(x), \vec{f}(x) \rangle$$

$$2. g(x_i) = x_i \wedge f(x_i) = p_i \Rightarrow E(X) = \sum_{i=1}^n p_i \cdot x_i = \mu$$

$$3. g(x_i) = x_i \wedge f(x_i) = p_i = \frac{1}{n} \Rightarrow E(X) = \frac{1}{n} \sum_{i=1}^n x_i = \mu$$

Wegen der Linearität von Summe und Integral gilt: • *A cause de la linéarité de la somme et de l'intégrale il vaut:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$g(X) = a h(X) + b u(X)$$

Beh.: • **Thè.:**

$$E(g(X)) = a E(h(X)) + b E(u(X))$$

4.7.4 Symmetrische Verteilung — Distribution symétrique

Definition: • **Définition:**

Eine Verteilung heisst **symmetrisch** bezüglich c , wenn gilt:

• *Une distribution s'appelle **symétrique** par rapport à c , s'il est:*

$$\forall_{c \pm x \in D_f} : f(c+x) = f(c-x)$$

Trivialerweise gilt für eine symmetrische Verteilung: • *Pour des raisons bien visibles, il vaut pour une distribution symétrique:*

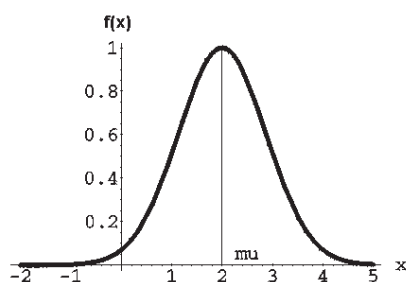
Satz: • Théorème: Vor.: • Hyp.:

f symmetrisch bezüglich c
• f symétrique par rapport à c

Beh.: • Thè.:

$\mu = c$

Bsp.: • Exemple:



4.7.5 Varianz, Standardabweichung — Variance, écart-type

Statt nach dem Erwartungswert von x resp. x_i zu fragen, fragen wir nach dem Erwartungswert der quadratischen Abweichung von μ von x resp. x_i . (Achtung: Der Erwartungswert der linearen Abweichung gibt immer 0. Diese Grösse macht daher keinen Sinn.)

• *Au lieu de s'intéresser à la valeur d'espérance de x resp. de x_i , nous considérons la valeur d'espérance de l'écart carré de μ de x resp. x_i . (Attention: La valeur d'espérance de l'écart linéaire donne toujours 0. Cette quantité par conséquent n'a pas de sens.)*

1. Diskreter Fall: • Cas discret:

Definition: • Définition:

Varianz • Variance (dispersion)

$$\sigma^2 = \sum_i (x_i - \mu)^2 \cdot f(x_i) \cdot 1$$

2. Stetiger Fall: • Cas continu:

Wie beim Erwartungswert wird im stetigen Fall auch hier die Summe zu einem Integral.

• *Comme à la valeur d'espérance aussi dans le cas continu la somme devient une intégrale.*

Definition: • Définition:

Varianz • Variance (dispersion)

$$\sigma^2 = \int_{-\infty}^{\infty} (x_i - \mu)^2 \cdot f(x_i) dx$$

3. Definition: • Définition:

$\sigma = \sqrt{\sigma^2}$ heisst **Standardabweichung**

• $\sigma = \sqrt{\sigma^2}$ s'appelle **écart type** (étalon, quadratique moyen), **écart-type** ou **déviatation normale**

Bemerkung: • **Remarque:** σ^2 ist wie erwähnt der Mittelwert der quadratischen Abweichung von μ . σ^2 gehört qualitativ als Mittelwert von Quadraten daher ebenfalls zur Klasse der Quadrate. σ ist somit ebenfalls ein Mass für die Grösse der mittleren Abweichung von μ , das qualitativ somit wieder zur Klasse der ursprünglichen nichtquadratischen Werte gehört. σ^2 oder auch σ werden manchmal als **Streuung** bezeichnet.

• σ^2 est, comme mentionné, la moyenne de l'écart carré de μ . Comme la moyenne de carrés, σ^2 appartient par conséquent aussi qualitativement à la classe des carrés. σ est donc par conséquent aussi une mesure pour le montant de l'écart moyen de μ , qui appartient par conséquent qualitativement encore à la classe des valeurs non-carrées et originales. σ^2 ou bien aussi σ sont parfois qualifiés de **dispersion**.

Konsequenz: • **Conséquence:** $g(X) = (X - \mu)^2 \Rightarrow E(g(X)) = E((X - \mu)^2) = \sigma^2$

Speziell: • *Spécialement:* (Diskreter Fall) • *(Cas discret)*

$$\begin{aligned} i \in \{1, 2, \dots, n\}, \quad g(x_i) = x_i, \quad f(x_i) = p_i \\ \Rightarrow \sigma^2 = E((X - \mu)^2) = E(X^2 - 2X\mu + \mu^2) = E(X^2) - E(2X\mu) + E(\mu^2) = E(X^2) - 2\mu E(X) + \mu^2 E(1) \\ = E(X^2) - 2\mu \cdot \mu + \mu^2 \cdot \frac{1}{n} \sum_{i=1}^n 1 = E(X^2) - 2\mu^2 + \mu^2 \cdot \frac{1}{n} \cdot n = E(X^2) - \mu^2 = E(X^2) - (E(X))^2 \end{aligned}$$

Korollar: • **Corollaire:** $i \in \{1, 2, \dots, n\}, \quad g(x_i) = x_i, \quad f(x_i) = p_i$

$$\Rightarrow \sigma^2 = E((X - \mu)^2) = E(X^2) - (E(X))^2$$

Bsp.: • **Exemple:**

X = „Anzahl Kopf“ beim einmaligen Wurf einer Münze.

• X = „Nombre de faces“ en jettant une monnaie une fois.

$$X = 0 \rightsquigarrow P(X = 0) = \frac{1}{2} \qquad X = 1 \rightsquigarrow P(X = 1) = \frac{1}{2}$$

$$\Rightarrow \mu = 0 \cdot \frac{1}{2} + 1 \cdot \frac{1}{2} = \frac{1}{2}$$

$$\Rightarrow \sigma^2 = (0 - \frac{1}{2})^2 \cdot \frac{1}{2} + (1 - \frac{1}{2})^2 \cdot \frac{1}{2} = (\frac{1}{2})^2 \cdot \frac{1}{2} + (\frac{1}{2})^2 \cdot \frac{1}{2} = \frac{1}{4}$$

$$\Rightarrow \sigma = \sqrt{\frac{1}{4}} = \frac{1}{2} \notin \{0, 1\}$$

Konsequenz: • **Conséquence:** $\sigma \in \{\dots, x_i, \dots\}$ ist nicht zwingend

• $\sigma \in \{\dots, x_i, \dots\}$ n'est pas toujours le cas.

4.7.6 Momente einer Verteilung — Moments d'une distribution

Grössen wie Mittelwert und Varianz heissen allgemein **Momente** der Verteilung der Zufallsvariablen X mit der Wahrscheinlichkeitsfunktion $f(X)$.

• Les grandeurs comme la moyenne et la variance s'appellent généralement des **moments** de la

distribution de la variable aléatoire X avec la fonction aléatoire $f(X)$.

Definition: • **Définition:**

Der folgende Erwartungswert heisst k -tes **Moment** der Verteilung:

• *La valeur d'espérance suivante s'appelle k -ème moment de la distribution:*

1. **Diskreter Fall:** • **Cas discret:**

$$m_k = E(X^k) = \sum_i x_i^k \cdot f(x_i)$$

$$m_{k,g} = E(g(X)^k) = \sum_i g(x_i)^k \cdot f(x_i)$$

2. **Stetiger Fall:** • **Cas continu:**

$$m_k = E(X^k) = \int_{-\infty}^{\infty} x^k \cdot f(x) dx$$

$$m_{k,g} = E(g(X)^k) = \int_{-\infty}^{\infty} g(x)^k \cdot f(x) dx$$

Idee: • **Idée:** $Y = X - \mu, f(x_i) = p = \frac{1}{n} \Rightarrow E(Y^k) = \sum_{i=1}^n \frac{(x_i - \mu)^k}{n} = \frac{1}{n} \cdot \sum_{i=1}^n (x_i - \mu)^k$
 $\leadsto$ Mechanik! • *Mécanique!*

Bsp.: • **Exemple:**

Mittelwert: • *Valeur moyenne:* $\mu = E(X) = E(X^1) = \sum_i x_i^1 \cdot f(x_i) = \sum_i x_i \cdot f(x_i)$

Varianz: • *Variance:* $\sigma^2 = E((X - \mu)^2) = \sum_i (x_i - \mu)^2 \cdot f(x_i)$

Satz: • **Théorème:**

$$1. E(X^2) = \sigma^2 + \mu^2 = E((X - E(X))^2) + (E(X))^2$$

$$2. E((X - \mu)^3) = E(X^3) - 3\mu E(X^2) + 2\mu^3$$

Beweis: • **Preuve:**

$$1. \sigma^2 = \sum_i (x_i - \mu)^2 \cdot f(x_i) = \sum_i x_i^2 \cdot f(x_i) - 2\mu \sum_i x_i \cdot f(x_i) + \mu^2 \cdot \sum_i f(x_i) = E(X^2) - 2\mu E(X) + \mu^2 \cdot 1 \\ = E(X^2) - \mu^2$$

2. Übung. • *Exercice.*

Konsequenz: • **Conséquence:** Daraus ergibt sich eine vereinfachte Berechnung der Varianz:

• *Il en résulte une façon simple de calculer la variance:*

Formel: • **Formule:** $\sigma^2 = E(X^2) - \mu^2 = E(X^2) - (E(X))^2, E(X^2) = \sum_i x_i^2 \cdot f(x_i)$

Schreibweise: • **Façon d'écrire:** Oft schreibt man auch: • *Souvent on trouve aussi:*

$$\mu_n =: E(X^n) = \begin{cases} \sum_i x_i^n p_i & \text{diskr.} \bullet \text{ discr.} \\ \int_{-\infty}^{\infty} x^n f(x) dx & \text{stet.} \bullet \text{ cont.} \end{cases}$$

$$\leadsto \text{Formel:} \bullet \text{ Formule:} \quad \sigma^2 = \mu_2 - \mu^2$$

Aus der Linearität der Summe und des Integrals ergibt sich:

• *De la linéarité de la somme et l'intégrale on déduit:*

$$\text{Satz:} \bullet \text{ Théorème:} \quad E(a_n X^n + a_{n-1} X^{n-1} + \dots + a_1 X + a_0) = a_n E(X^n) + a_{n-1} E(X^{n-1}) + \dots + a_1 E(X) + a_0 = a_n \mu_n + a_{n-1} \mu_{n-1} + \dots + a_1 \mu_1 + a_0$$

Daraus leiten wir ab: • *A l'aide de cette formule on déduit:*

$$\begin{aligned} E(aX + b) &= aE(X) + b \Rightarrow \mu(aX + b) = a\mu(X) + b \Rightarrow \\ (\sigma(aX + b))^2 &= E((aX + b - \mu(aX + b))^2) = E((aX + b)^2 - 2(aX + b)(\mu(aX + b)) + (\mu(aX + b))^2) = \\ &= E(a^2 X^2 + 2aXb + b^2 - 2(aX + b)(a\mu(X) + b) + (a\mu(X) + b)^2) = \\ &= E(a^2 X^2 + 2aXb + b^2 - 2a^2 X\mu(X) - 2aXb - 2ba\mu(X) - 2b^2 + a^2 \mu(X)^2 + 2a\mu(X)b + b^2) = \\ &= E(a^2 X^2 - 2a^2 X\mu(X) + a^2 \mu(X)^2) = E(a^2 (X - \mu(X))^2) = a^2 E((X - \mu(X))^2) = a^2 \sigma^2 \leadsto \end{aligned}$$

$$\text{Satz:} \bullet \text{ Théorème:} \quad (\sigma(aX + b))^2 = a^2 \sigma(X)^2 \\ \text{oder} \bullet \text{ ou } \text{Var}(aX + b) = a^2 \text{Var}(X)$$

Der letzte Satz erlaubt manchmal eine Vereinfachung der Berechnung der Varianz.

• *Le théorème ci-dessus permet parfois une simplification du calcul de la variance.*

$$\text{Korollar:} \bullet \text{ Corollaire:} \quad (\sigma(b))^2 = \text{Var}(b) = 0$$

Man rechnet sofort nach: • *On peut vérifier tout de suite par le calcul:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Zu X gehöre der Mittelwert μ und die Varianz σ^2 .

• *Pour X on a la moyenne μ et la variance σ^2 .*

$$Z = \frac{X - \mu}{\sigma}$$

Beh.: • **Thè.:**

Zu Z gehöre der Mittelwert 0 und die Varianz 1.

• *Pour Z on a la moyenne 0 et la variance 1.*

Zum Beweis: • **Quant à la preuve:**

$$\mu(Z) = E(Z) = \frac{1}{\sigma} \cdot (E(X) - E(\mu)) = \frac{1}{\sigma} \cdot (\mu - \mu \cdot E(1)) = \frac{1}{\sigma} \cdot (\mu - \mu \cdot 1) = 0$$

$$\begin{aligned} \text{P. 98: } E(X^2) = \sigma^2 + \mu^2 &\Rightarrow \sigma(Z)^2 = E(Z^2) - \mu(Z)^2 = E(Z^2) = E\left(\frac{1}{\sigma^2} \cdot (X^2 - 2X\mu + \mu^2)\right) \\ &= \frac{1}{\sigma^2} \cdot (E(X^2) - 2E(X)\mu + \mu^2) = \frac{1}{\sigma^2} \cdot (E(X^2) - \mu^2) = \frac{1}{\sigma^2} \cdot (\sigma^2 + \mu^2 - \mu^2) = 1 \end{aligned}$$

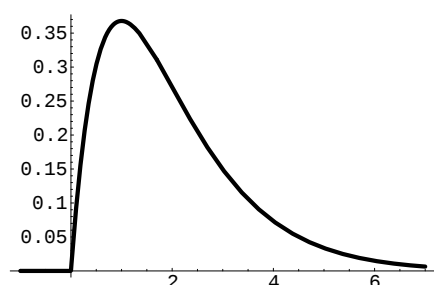
Infolge des letzten Satzes definieren wir: • *Suivant le théorème ci-dessus nous définissons:*

Definition: • **Définition:** $Z = \frac{X - \mu}{\frac{\sigma}{\sigma}}$ heisst **Standardform** von X .
 • $Z = \frac{X - \mu}{\sigma}$ s'appelle **forme standard** de X .

4.7.7 Schiefe einer Verteilung — Dissymétrie d'une distribution

Definition: • **Définition:** Sei • *Soit* $\gamma = \frac{1}{\sigma^3} \cdot E((X - \mu)^3)$
 γ heisst **Schiefe** der Verteilung. • γ s'appelle **dissymétrie** de la distribution.

Bemerkung: • **Remarque:** Die Schiefe der Verteilung ist ein Mass für die Assymetrie der Verteilung.
 • *La dissymétrie est une mesure pour le manque de symétrie de la distribution.*



$$f(x) = \begin{cases} 0 & x < 0 \\ x \cdot e^{-x} & x \geq 0 \end{cases}$$

$$\leadsto \int_{-\infty}^{\infty} f(x) dx = 1, \quad \gamma = ?$$

Mit partieller Integration erhalten wir: • *A l'aide de l'intégration partielle nous recevons:*

$$\begin{aligned} \mu = E(X) &= \int_0^{\infty} x \cdot e^{-x} \cdot x dx = \int_0^{\infty} x^2 \cdot e^{-x} dx = 2, \quad E(X^2) = 6, \quad E(X^3) = 24 \\ \Rightarrow E((X - \mu)^3) &= 24 - 3 \cdot 2 \cdot 6 + 2 \cdot 2^3 = 4, \quad \sigma^2 = E(X^2) - (E(X))^2 = 6 - 2^2 = 2 \\ &\Rightarrow \gamma = \frac{1}{\sigma^3} \cdot E((X - \mu)^3) = \frac{1}{2(\frac{3}{2})} \cdot 4 = \sqrt{2} \approx 1.41421 \end{aligned}$$

4.7.8 Weitere Kenngrössen — Autres caractéristiques

Variationskoeffizient — Coefficient de variation

Der **Variationskoeffizient** oder **Variabilitätskoeffizient** ν einer Zufallsgrösse X mit $\mu \neq 0$ ist wie folgt definiert: • *Le coefficient de variation ou coefficient de variabilité ν d'une variable aléatoire X avec $\mu \neq 0$ est défini comme il suit:*

Definition: • **Définition:** $\nu = \frac{\sigma}{\mu}$

Median, Modus — Médian, mode

Median und Modus sind Lagemasszahlen. Sie betreffen vor allem die beurteilende Statistik. (Vgl. 2.)

• *Médian et mode sont des mesures pour la position du graphe d'une fonction. Ils concernent en particulier la statistique descriptive. (Voir 2.)*

Definition: • Définition:

Der **Median** einer Menge von Zahlen, die der Grösse nach geordnet sind, ist der Wert in der Mitte oder das arithmetische Mittel der beiden Werte in der Mitte der geordneten Datenfolge.

• *Le médian d'un ensemble de nombres qui sont ordonnés selon la grandeur, est la valeur du milieu ou la moyenne arithmétique des deux valeurs du milieu de la suite rangée des données.*

Falls die Daten in Klassen eingeteilt worden sind, muss die Definition angepasst werden (vgl. Lit.).

• *Si les données ont été groupées en classes, la définition doit être adaptée à cette situation, voir. lit.*

Quantile, Fraktile — Quantiles, fractiles

Gegeben sei eine stetige Zufallsgrösse X mit einer Dichtefunktion f und einer Verteilungsfunktion F . Sei $0 < q < 1$.

• *Soit donné une variable aléatoire continue X avec une fonction de densité f et une fonction de distribution F . Soit $0 < q < 1$.*

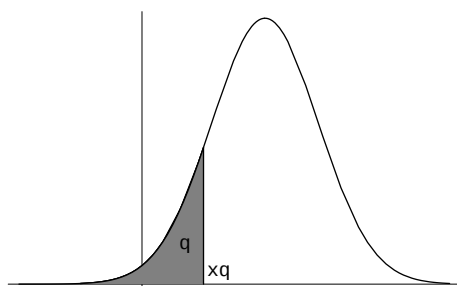
Definition: • Définition:

$x_q \in \mathbb{R}$ heisst das **untere Quantil der Ordnung q (q -Quantil)**, wenn gilt: $F(x_q) = P(X \leq x_q) = q$
Entsprechend für das **obere Quantil**.

• $x_q \in \mathbb{R}$ s'appelle **quantile inférieur de l'ordre q (quantile q)**, s'il vaut: $F(x_q) = P(X \leq x_q) = q$
Pareillement pour le **quantile supérieur**.

Konsequenz: • Conséquence: Das Quantil der Ordnung $q = 0.5$ ist gerade der Median.

• *Le quantile de l'ordre $q = 0.5$ est exactement le médian.*



Spezielle Quantile sind **Quantile**

($Q_1 = x_{0.25}$, $Q_3 = x_{0.75}$), **Quintile**, **Dezile** oder **Zentile** (P_{10}, P_{90}) u.s.w..

• *Voici des quantiles spéciaux: Les quantiles ($Q_1 = x_{0.25}$, $Q_3 = x_{0.75}$), les quintiles, les deciles ou les centiles (P_{10}, P_{90}) etc..*

Bemerkung: • Remarque:

Der **Quartilabstand** $Q_3 - Q_1 = x_{0.75} - x_{0.25}$ ist ein Streuungsmass. Ebenso der Zentilabstand $P_{90} - P_{10}$. Oft wird auch der **halbe Quartilabstand** $Q = \frac{Q_3 - Q_1}{2}$ benutzt.

• *La distance des quartiles $Q_3 - Q_1 = x_{0.75} - x_{0.25}$ est une mesure de dispersion. De même la distance des centiles $P_{90} - P_{10}$. Souvent on utilise aussi la demi-distance des quartiles*

$$Q = \frac{Q_3 - Q_1}{2}.$$

Exzess, Wölbung — Excès, kurtosis, courbure, convexité

Z.B. der **Exzess (Kurtosis)** κ ist ein Mass für die Steilheit einer Verteilung. • *P.ex. le kurtosis κ est une mesure pour la pente d'une distribution.*

Definition: • **Définition:**

Exzess (Kurtosis): • **kurtosis:**

$$\kappa = \frac{Q}{P_{90} - P_{10}} \quad (Q, P \text{ vgl. oben} \quad \bullet \text{ voir en haut})$$

Momentenkoeffizient der Kurtosis: • **Coefficient de moments du kurtosis:**

$$a_4 = \frac{m_4}{m_2^2} \quad (m_k = E(X^k))$$

Andere Mittelwerte — Autres moyennes

Für andere Mittelwerte wie **harmonische** oder **geometrische** Mittelwerte sowie **empirische** Beziehungen vgl. Lit..

• *Pour les autres moyennes comme les moyennes harmoniques ou géométriques ainsi que relations empiriques voir lit..*

4.7.9 Momentenerzeugende charakteristische Funktion — Fonction caractéristique génératrice de moments

Der Grund zu der folgenden Definition wird klar, wenn wir damit arbeiten:

• *En travaillant avec la définition suivante, on comprend la raison pour laquelle on l'utilise:*

Definition: • **Définition:**

Die Erwartung $E(e^{tX})$ von e^{tX} heisst **momentenerzeugende Funktion** $G(t)$.

• *La valeur espérée $E(e^{tx})$ de e^{tx} s'appelle fonction caractéristique génératrice de moments $G(t)$.*

Wir berechnen $G(t)$ für beliebige Funktionen $f(x)$:

• *Nous calculons $G(t)$ pour des fonctions quelconques $f(x)$:*

Diskreter Fall: • **Cas discret:** $G(t) = E(e^{tX}) = \sum_i e^{tx_i} \cdot f(x_i)$

Für die k -te Ableitung nach t erhalten wir: • *Pour la k -ème dérivée d'après t nous obtenons:*

$$\frac{d^k G(t)}{dt^k} = G_t^{(k)}(t) = \sum_i x_i^k \cdot e^{tx_i} \cdot f(x_i)$$

Kontinuierlicher Fall: • **Cas continu:** $G(t) = E(e^{tx}) = \int_{-\infty}^{\infty} e^{tx} \cdot f(x) dx$

$$\frac{d^k G(t)}{dt^k} = G_t^{(k)}(t) = \int_{-\infty}^{\infty} x^k \cdot e^{tx} \cdot f(x) dx$$

Satz: • **Théorème:** $\leadsto$ **Trick:** • **Truc:**

$$t = 0 \Rightarrow G^{(k)}(0) = E(X^k) = \begin{cases} \sum_i x_i^k \cdot e^{0 \cdot x_i} \cdot f(x_i) & = \sum_i x_i^k \cdot f(x_i) \\ \int_{-\infty}^{\infty} x^k \cdot e^{0 \cdot x} \cdot f(x) dx & = \int_{-\infty}^{\infty} x^k \cdot f(x) dx \end{cases}$$

Speziell: • **Spécialement:**

$$k = 1 \Rightarrow G'_t(0) = E(X^1) = E(X) = \mu \rightsquigarrow \text{Mittelwert} \bullet \text{Valeur moyenne}$$

$$\Rightarrow \sigma^2 = E(X^2) - \mu^2 = G''_t(0) - G'_t(0)^2 \rightsquigarrow \text{Varianz} \bullet \text{Variance}$$

Bemerkung: • **Remarque:**

Die momenterzeugende Funktion wird uns im nächsten Kapitel bei der Bernoulli-Verteilung nützlich sein.

• *La fonction génératrice de moments nous sera utile dans le chapitre prochain au sujet de la distribution de Bernoulli.*

4.7.10 Laplace- und z-Transformation — Transf. de Laplace et en z

Dem mit der Ingenieurmathematik vertrauten Leser fällt auf, dass die momenterzeugende Funktion $G(t)$ offensichtlich mit der Laplace- resp. der z -Transformation zu tun hat. Der Zusammenhang wird bei späterer Gelegenheit genauer erläutert, da er im Moment nicht unverzichtbar ist. Kenntnisse der Laplace- und der z Transformation erwirbt man in der höheren Analysis, vgl. Lit. Fortsetzung Mathematik, Wirz, Bibl. A16.

• *Le lecteur qui connaît les mathématiques d'ingénieur remarque que la fonction génératrice de moments $G(t)$ évidemment à faire avec les transformations de Laplace resp. les transformations en z . Le rapport est expliqué plus exactement plus tard parce qu'il n'est pas indispensable en ce moment. On acquiert les connaissances de la transformation de Laplace et de la transformation en z dans l'analyse supérieure, voir lit. Suite mathématiques, Wirz, Bibl. A16.*

4.8 Spezielle diskrete Verteilungen — Distributions discrètes spéciales

4.8.1 Bernoulliverteilung — Distribution de Bernoulli

Situation und Einzelversuch — **Situation et expérience indépendante**

Problem: • **Problème:**

Wir studieren die Anzahl des Eintreffens eines bestimmten Ereignisses A bei der wiederholten Ausführung eines Zufallsexperiments, bei dem die Ereignisse unabhängig und gleichwahrscheinlich sind. Z.B. sei A das Eintreffen einer 2 beim Würfeln mit einem Würfel, wobei mehrmals gewürfelt wird. Das betrachtete Ereignis kann eintreffen oder nicht. Die dazugehörige Zufallsvariable X kann daher nur zwei Werte annehmen: 0 und 1.

• *Nous étudions le nombre de réalisations d'un certain événement A qui soit exécuté répétitivement pendant une expérience aléatoire dont les événements sont indépendants et ont la même probabilité. Soit A p.ex. la réalisation d'un 2 en jouant aux dés avec un dé, le dé sera lancé plusieurs fois. L'événement considéré peut être réalisé ou ne pas être réalisé. La variable aléatoire X nécessaire peut donc atteindre seulement deux valeurs: 0 et 1.*

Sei $X = 0$ falls A nicht eintritt und $X = 1$ falls A eintritt.

• *Soit $X = 0$ si A n'est pas réalisé et $X = 1$ si A est réalisé.*

Sei • *Soit* p = Wahrscheinlichkeit dass A eintritt. • p = probabilité que A soit réalisé.
 $q = 1 - p$ = Wahrscheinlichkeit dass A nicht eintritt.
 • $q = 1 - p$ = probabilité que A ne soit pas réalisé.
 X = Zufallsvariable: Anzahl des Eintreffens von A .
 • X = Variable aléatoire: Nombre de réalisations de A .

Folgerung: • **Conclusion:** Erfolgswahrscheinlichkeit beim Einzelversuch:
 • Probabilité de succès à l'expérience sans répétition:
 $P(X = 0) = q, f(0) = q$
 $P(X = 1) = p, f(1) = p$
 $x \notin \{0, 1\} \Rightarrow f(x) = 0$

Dieser Sachverhalt kann kurz in einer einzigen Formel wie folgt gefasst werden:

• Ce fait peut être écrit brièvement dans une seule formule:

Konsequenz: • **Conséquence:** Im Einzelversuch gilt: • Lors d'une expérience unique il vaut:

$$f(x) = \begin{cases} p^x \cdot q^{1-x} & x \in \{0, 1\} \\ 0 & x \notin \{0, 1\} \end{cases}$$

Definition: • **Définition:** Den eben beschriebenen Versuch nennt man oft auch **Bernoulli-Versuch**. • L'expérience qu'on vient de décrire s'appelle souvent expérience de Bernoulli.

Mehrfachversuch — Expérience exécutée plusieurs fois

Bei n -maliger Wiederholung des Experiments kann sich der Erfolg $0, 1, 2, \dots$ oder n mal einstellen. Erfolg heisst, dass A eintritt. Daher kann also X die Werte $0, 1, 2, \dots$ oder n annehmen.

• Si l'expérience est répétée n -fois le succès peut arriver $0, 1, 2, \dots$ ou n fois. Succès veut dire que A est réalisé. X peut donc prendre les valeurs $0, 1, 2, \dots$ ou n .

Bei n Versuchen treffe nun A x mal ein. Da die Versuche identisch und daher die Ereignisse unabhängig sind, ist die Wahrscheinlichkeit des Gesamtereignisses das Produkt der Wahrscheinlichkeiten der Einzelereignisse:

• L'événement A soit réalisé x fois à n répétitions de l'expérience. Comme les expériences partielles sont identiques et par conséquent indépendantes, la probabilité de l'événement total est le produit des probabilités des événements partiels:

$$P(X = x) = \underbrace{p \cdot p \cdot \dots \cdot p}_{p^x} \cdot \underbrace{q \cdot q \cdot \dots \cdot q}_{q^{n-x}} = p^x \cdot q^{n-x}$$

Konsequenz: • **Conséquence:** Die Formel für $n = 1$ ist hier eingeschlossen.

• La formule pour $n = 1$ est donc encore valable.

Der Erfolg kann sich bei gleicher Gesamtzahl auf verschiedene Arten auf die Einzelereignisse verteilen. Die Anzahl Möglichkeiten erhält man aus der Antwort auf die Frage, auf wieviele Arten man aus n Elementen x Elemente ohne Wiederholung auswählen kann. Das führt daher auf eine Kombination ohne Wiederholung:

• Au même nombre total, le succès peut être réalisé de manières différentes par les événements partiels. On obtient le nombre de possibilités de la réponse à la question suivante: De combien de manières différentes

est-ce qu'on peut choisir x éléments parmi n éléments sans répétition. Ça nous mène par conséquent à une combinaison sans répétition:

$$\text{Möglichkeiten: } \bullet \text{ Possibilités: } \binom{n}{x} = \frac{n!}{x!(n-x)!}$$

Bsp.: • Exemple: Bei 3 Versuchen soll ein Ereignis 2 mal eintreten. Wieviele Möglichkeiten gibt es? — Der Erfolg kann sich (a) im 1. und 2. Versuch oder (b) im 1. und 3. Versuch oder (c) im 2. und 3. Versuch einstellen. Es gibt also $|\{(a), (b), (c)\}| = 3 = \binom{3}{2}$ Möglichkeiten.

• *Lors de 3 expériences un événement doit se réaliser 2 fois. Combien de possibilités est-ce qu'il y a? — Le succès peut se réaliser (a) lors des 1. et 2. expériences ou (b) lors des 1. et 3. expériences ou (c) lors des 2. et 3. expériences. Il y a donc $|\{(a), (b), (c)\}| = 3 = \binom{3}{2}$ possibilités.*

Satz: • Théorème: **Vor.: • Hyp.:**

Bernoulli-Experiment mit n Wiederholungen
• *Expérience de Bernoulli avec n répétitions*

Beh.: • Thè.:

$$f(x) = P(X = x) = \binom{n}{x} p^x q^{n-x}$$

Definition: • Définition: Die durch $f(x) = \binom{n}{x} p^x q^{n-x}$ gegebene Verteilung mit den Parametern $n \in \mathbb{N}$ und $p \in (0, 1)$ heisst **Bernoulli-Verteilung** oder **Binomialverteilung** $Bi(n, p)$.

• *La distribution donnée par $f(x) = \binom{n}{x} p^x q^{n-x}$ avec les paramètres $n \in \mathbb{N}$ et $p \in (0, 1)$ s'appelle **distribution de Bernoulli** oder **distribution binomiale** $Bi(n, p)$.*

Korollar: • Corollaire: $1 = \sum_{x=0}^n P(X = x) = \sum_{x=0}^n \binom{n}{x} p^x q^{n-x}$

Bemerkung: • Remarque: Aus $p + q = 1$ bei n unabhängigen Wiederholungen deduziert man ebenfalls:

• *De $p + q = 1$ à n répétitions indépendantes on déduit aussi:*

$$1 = 1^n = (p + q)^n = \sum_{x=0}^n \binom{n}{x} p^x q^{n-x}$$

4.8.2 Gesetze für die Binomialverteilung — Lois pour la distribution de Bernoulli

Die Verteilungsfunktion ergibt sich durch eine einfache Summation:

• *La fonction de distribution est obtenue de façon simple en prenant la somme:*

Korollar: • Corollaire: $F(x) = \sum_{k \leq x} \binom{n}{k} p^k q^{n-k}, \quad x \geq 0$

Weiter gilt: • *En outre il vaut:* $\binom{n}{x} = \frac{n!}{x! \cdot (n-x)!}$

Daraus folgt: • *On en déduit:* $\binom{n}{x+1} = \frac{n-x}{x+1} \cdot \binom{n}{x}$

Damit erhalten wir eine Rekursionsformel: • *Ça nous donne une formule de récurrence:*

Korollar: • **Corollaire:** $f(x+1) = \left(\frac{n-x}{x+1}\right) \cdot \frac{p}{q} \cdot f(x)$

Bsp.: • **Exemple:**

Es wird eine Sendung mit N Schrauben vorbereitet. M davon sind unbrauchbar (Ausschuss). Die Sendung wird vom Kunden geprüft, indem er n mal eine Schraube zieht und wieder zurücklegt. Wie gross ist die Chance, dass der Kunde x unbrauchbare Schrauben findet?

• *On prépare une livraison avec N vis. M pièces sont inutilisables (rebut). La livraison est examinée par le client. Il en tire n fois une vis et la remet. Quelle est la chance que le client trouve x vis inutilisables?*

Offensichtlich handelt es sich um ein Bernoulliexperiment. Die Chance, bei einmaligem Ziehen Ausschuss zu finden ist $p = \frac{M}{N}$.

• *Il est bien visible qu'il s'agit d'une expérience de Bernoulli. La chance de trouver du rebut en tirant une seule fois est $p = \frac{M}{N}$.*

$$\leadsto f(x) = \binom{n}{x} \cdot \left(\frac{M}{N}\right)^x \cdot \left(1 - \frac{M}{N}\right)^{n-x}$$

Zur Berechnung von Mittelwert und Varianz benutzen wir die momentenerzeugende Funktion:

• *Pour calculer la moyenne et la variance nous utilisons la fonction génératrice de moments:*

$$G(t) = E(e^{tx}) = \sum_i e^{tx_i} \cdot f(x_i) = \sum_{x=0}^n e^{tx} \cdot \binom{n}{x} p^x q^{n-x} = \sum_{x=0}^n \binom{n}{x} \cdot (p \cdot e^t)^x \cdot q^{n-x} = (p \cdot e^t + q)^n$$

$$\Rightarrow G'(t) = n \cdot (p \cdot e^t + q)^{n-1} \cdot p \cdot e^t,$$

$$G''(t) = n \cdot (n-1) \cdot (p \cdot e^t + q)^{n-2} \cdot p^2 \cdot e^{2t} + n \cdot (p \cdot e^t + q)^{n-1} \cdot p \cdot e^t = n \cdot p \cdot e^t \cdot (p \cdot e^t + q)^{n-2} \cdot ((n-1) \cdot p \cdot e^t + (p \cdot e^t + q))$$

$$= e^t n p (e^t p + q)^{n-2} (e^t n p + q) \Rightarrow E(X^2) = G''(0) = n p \underbrace{(p + q)^{n-2}}_{=1} (n p + q) = n \cdot p (n p + q)$$

$$\Rightarrow \mu = G'(0) = n \cdot p, \quad \sigma^2 = G''(0) - \mu^2 = E(X^2) - (E(X))^2 = n^2 \cdot p^2 + n \cdot p \cdot q - n^2 \cdot p^2 = n \cdot p \cdot q$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Bernoulliverteilung

• *Distribution de Bernoulli*

Beh.: • **Thè.:**

$$\mu = n \cdot p$$

$$\sigma^2 = n \cdot p \cdot q$$

Bsp.: • **Exemple:**

Kartenspiel, 32 Karten mit 4 Assen. Experiment: 6 mal eine Karte ziehen mit zurücklegen (Bernoulliexperiment). Ereignis A : mindestens 3 Assen kommen.

• *Jeu de cartes, 32 cartes avec 4 as. L'expérience: Tierer 6 fois une carte et la remettre (expérience de Bernoulli). Événement A: Tirer au moins 3 as.*

$\leadsto$ Einzelexperiment: $p = \frac{4}{32}$ • *Expérience indépendante: $p = \frac{4}{32}$*

$\leadsto$ $A : 3, 4, 5$ oder 6 mal ein As. • $A : 3, 4, 5$ ou 6 fois un as.

Kurz: • *Bref:* $A = \{3, 4, 5, 6\} = \{3\} \cup \{4\} \cup \{5\} \cup \{6\}$

$$\Rightarrow P(A) = P(\{3\}) + P(\{4\}) + P(\{5\}) + P(\{6\}) = \sum_{k=3}^6 P(\{k\}) = \sum_{k=3}^6 \binom{6}{k} \left(\frac{4}{32}\right)^k \left(1 - \frac{4}{32}\right)^{6-k} \approx 0.03.$$

Zum Vergleich: • *Pour comparaison:* $\mu = \frac{3}{4} = 0.74$, $\sigma^2 = \frac{21}{32} \approx 0.656$, $\sigma \approx 0.81$

4.8.3 Poissonverteilung — Distribution de Poisson

Verteilungsfunktion — Fonction de distribution

Problem: • **Problème:** Gesucht ist eine Verteilung, die für grosse n und kleine p einerseits die Bernoulliverteilung gut annähert und andererseits mathematisch bequemer zu handhaben ist. Das Gewünschte leistet die **Poissonverteilung**.

• *On cherche une distribution qui d'une part approche relativement bien la distribution de Bernoulli pour les grands n et des petits p et qui d'autre part est commode pour des manipulations mathématiques. Ce qu'on veut est effectué par la **distribution de Poisson**.*

Idee: • **Idée:** Für die Bernoulliverteilung gilt: • *Pour la distribution de Bernoulli il vaut:* $\mu = n \cdot p$

$$\begin{aligned} \leadsto p &= \frac{\mu}{n}, \quad p^x = \frac{\mu^x}{n^x}, \quad q^{n-x} = (1-p)^{n-x} = \left(1 - \frac{\mu}{n}\right)^{n-x} = \left(1 - \frac{\mu}{n}\right)^n \cdot \left(1 - \frac{\mu}{n}\right)^{-x} \\ \Rightarrow f(x) &= \binom{n}{x} \cdot p^x \cdot q^{n-x} = \frac{n \cdot (n-1) \cdot \dots \cdot (n-x+1)}{x!} \cdot \frac{\mu^x}{n^x} \cdot \left(1 - \frac{\mu}{n}\right)^n \cdot \left(1 - \frac{\mu}{n}\right)^{-x} = \\ &= \frac{n \cdot (n-1) \cdot \dots \cdot (n-x+1)}{n^x} \cdot \left(1 - \frac{\mu}{n}\right)^{-x} \cdot \frac{\mu^x}{x!} \cdot \left(1 - \frac{\mu}{n}\right)^n \end{aligned}$$

Sei nun $n \rightarrow \infty$ und $p \rightarrow 0$ derart dass gilt $\mu = n \cdot p = \lim_{n \rightarrow \infty, p \rightarrow 0} n \cdot p$.

• *Soit maintenant $n \rightarrow \infty$ et $p \rightarrow 0$ de façon qu'il vaut $\mu = n \cdot p = \lim_{n \rightarrow \infty, p \rightarrow 0} n \cdot p$.*

$$\leadsto \frac{n \cdot (n-1) \cdot \dots \cdot (n-x+1)}{n^x} = \frac{n}{n} \cdot \frac{n-1}{n} \cdot \dots \cdot \frac{n-x+1}{n} \xrightarrow{n \rightarrow \infty, x \ll n} 1$$

Weiter gilt: • *En outre il vaut:* $\left(1 - \frac{\mu}{n}\right)^n = \left(1 + \frac{-\mu}{n}\right)^n \xrightarrow{n \rightarrow \infty} e^{-\mu}$, $\left(1 - \frac{\mu}{n}\right)^{-x} \xrightarrow{n \rightarrow \infty, x, \mu \ll n} 1$

Damit erhalten wir zusammengefasst das folgende Resultat (**Grenzwertsatz von Poisson**): • *Nous obtenons donc en somme le résultat suivant (Théorème limite de Poisson):*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\mu = n \cdot p = \lim_{n \rightarrow \infty, p \rightarrow 0} n \cdot p$$

Beh.: • **Thè.:**

$$f(x) \xrightarrow{n \rightarrow \infty, x, \mu \ll n} \frac{\mu^x}{x!} \cdot e^{-\mu}$$

Definition: • **Définition:** $f(x) = \frac{\mu^x}{x!} \cdot e^{-\mu}$ heisst **Poissonverteilung** $Po(\mu)$ *¹

• *s'appelle distribution de Poisson $Po(\mu)$*

Konsequenz: • **Conséquence:** $x \geq 0 \Rightarrow F(x) = e^{-\mu} \cdot \sum_{s \leq x} \frac{\mu^s}{s!}$

Bemerkung: • **Remarque:** Da wegen $p \rightarrow 0$ die Wahrscheinlichkeiten der untersuchten Ereignisse sehr klein werden können, spricht man auch von der Verteilung der **seltenen Ereignissen**.
 • *Comme par conséquent de $p \rightarrow 0$ les probabilités des évènements examinés peuvent être très petites, on parle aussi d'évènements rares.*

Momente — Moments

$$G(t) = \sum_i e^{tx_i} f(x_i) = \sum_i e^{tx_i} e^{-\mu} \cdot \frac{\mu^{x_i}}{x_i!} \Rightarrow G'(t) = e^{-\mu+e^t \mu} \mu \Rightarrow G'(0) = \mu \text{ O.k.!} \bullet \text{ O.k.!}$$

Damit ist aber für die Praxis nichts gewonnen. Um weiterzukommen erinnern wir uns, dass die Poissonverteilung eine Annäherung der Binomialverteilung für $x \ll n$ und $n \rightarrow \infty$ ist. Daher können wir in der Praxis für μ den Mittelwert der beobachteten relativen Häufigkeiten nehmen, wenn die genannten Bedingungen annähernd erfüllt sind und es sich um ein Bernoulliexperiment handelt.

• *Par ce calcul on n'a rien gagné pour la pratique. Pour avancer, nous nous souvenons que la distribution de Poisson est une approximation de la distribution binomiale pour $x \ll n$ et $n \rightarrow \infty$. Par conséquent nous pouvons prendre la moyenne des fréquences relatives et observées dans la pratique pour la moyenne μ si les conditions données sont satisfaites et si s'il s'agit d'une expérience de Bernoulli.*

$$\leadsto \mu \approx \frac{1}{n} \cdot \sum_{k=1}^j n_k \cdot x_k$$

$$G''(t) = e^{-\mu+e^t \mu} \mu (1 + e^t \mu) \Rightarrow G''(0) = \mu(1 + \mu) = \sigma^2 + \mu^2 \Rightarrow \sigma^2 = \mu^2, \sigma = \mu$$

$$\text{Ebenso berechnet man: } \bullet \text{ Analogiquement on calcule: } \gamma = \frac{1}{\sigma^3} \cdot E((X - \mu)^3) = \frac{1}{\sqrt{\mu}}$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Poissonverteilung • *Distribution de Poisson*

Beh.: • **Thè.:**

$$\begin{aligned} \mu &\approx \frac{1}{n} \cdot \sum_{k=1}^j n_k \cdot x_k \\ \sigma^2 &= \mu^2, \sigma = \mu \\ \gamma &= \frac{1}{\sqrt{\mu}} \end{aligned}$$

Bsp.: • **Exemple:**

Wie gross ist die Wahrscheinlichkeit, dass an einer Schule mit $n = 1000$ Studierenden mindestens $k = 1$ Studierende am 1. Mai Geburtstag haben?

• *Quelle est la probabilité qu'à une école avec $n = 1000$ étudiants il y a au moins $k = 1$ étudiants qui ont leur anniversaire le 1er mai?*

Bernoulliexperiment: • *Expérience de Bernoulli:*

$$p = \frac{1}{365}, n = 1000, P(X \geq 1) = 1 - P(X < 1) = 1 - P(X = 0) = 1 - \left(\frac{365-1}{365}\right)^{1000} \approx 0.935654.$$

Poissonapproximation: • *Approximation de Poisson:*

$$P(X = k) \approx \frac{\lambda^k}{k!} \cdot e^{-\lambda}, \quad \lambda = n \cdot p = \frac{1000}{365}, \quad P(X \geq 1) = 1 - P(X < 1) = 1 - P(X = 0) \approx 1 - \frac{\lambda^0}{0!} \cdot e^{-\lambda} \approx 0.935412.$$

4.8.4 Pascalverteilung — Distribution de Pascal

Das Experiment bestehe wieder aus unabhängigen Versuchen. Sei p im Gegensatz zum Bernoulliexperiment die gegebene Wahrscheinlichkeit für einen **Misserfolg** im Basisversuch. $1 - p$ ist daher hier die Wahrscheinlichkeit für den Erfolg. k sei der Parameter der Erfolge bis zum r -ten **Misserfolg** (**Schaden**). X variiert auf $\{k_i\}$. An die Stelle von p^x bei der Bernoulliverteilung tritt jetzt $(1 - p)^k$, an die Stelle von $(1 - p)^{n-x}$ tritt p^r ($r = n - x = n - k$). Ausgewählt werden k Möglichkeiten aus einer Menge von $r + k - 1$ Möglichkeiten. (Die r Misserfolge beim Experiment ohne den letzten, der immer an letzter Stelle kommt und das Experiment abbricht, also die Anzahl der Auswahlmöglichkeiten nicht erhöht. k ist die Anzahl Erfolge.) Die Lösung des Auswahlproblems führt dann zu folgender Formel:

• *L'expérience consiste de nouveau en expériences partielles et indépendantes. Soit p par contre à l'expérience de Bernoulli la probabilité donnée d'un échec dans l'expérience de base. $1 - p$ est donc par conséquent la probabilité pour le succès. k soit le paramètre des succès jusqu'au r -ème échec (ne pas réussir). X varie sur $\{k_i\}$. $(1 - p)^k$ prend maintenant la place de p^x dans la distribution de Bernoulli, p^r prend la place de $(1 - p)^{n-x}$ ($r = n - x = n - k$). (On choisit k possibilités dans un ensemble de $r + k - 1$ possibilités, parce que lors de l'expérience on a r échecs sans le dernier qui vient toujours à la dernière place et termine donc l'expérience, ce qui n'augmente pas les possibilités du choix. k est le nombre de succès.) La solution du problème de choix mène donc à la formule suivante:*

Formel: • **Formule:** $X =$ Anzahl unabhängiger Versuche bis zum r -ten Misserfolg, $q =$ Erfolgswahrscheinlichkeit (p : Misserfolg), $n := k + r$. • $X =$ nombre d'expérience indépendantes jusqu'à ne pas réussir pour le r -ème fois, $q =$ probabilité de succès (p : échec), $n := k + r$.

$$nB_k := p_k = P(X = k) = \binom{r+k-1}{k} \cdot p^r \cdot (1-p)^k, \quad k = 0, 1, \dots, \quad P(X = n) = \binom{n-1}{r-1} \cdot p^r \cdot (1-p)^{n-r}$$

Da hier der Misserfolg p bestimmt, ist die folgende Definition plausibel: • *Comme ici l'échec détermine p , la définition suivante devient compréhensible:*

Definition: • **Définition:** $r > 0$, $0 < p < 1 \rightsquigarrow nB_k$ heisst **negative Binomialverteilung**.
• $r > 0$, $0 < p < 1 \rightsquigarrow nB_k$ s'appelle **distribution binomiale négative**.

Definition: • **Définition:** Für $r \in \mathbb{N}$ heisst die Verteilung nB_k auch **Pascal-Verteilung** $Pa(r, p)$.
• Pour $r \in \mathbb{N}$ la distribution nB_k s'appelle aussi **distribution de Pascal** $Pa(r, p)$.

Anwendungen finden sich in der Schadenstheorie.

• *On trouve des applications dans la théorie du dégât.*

Die direkte Berechnung von μ und σ^2 von Hand ist etwas mühsam. Mit der bei der Binomialverteilung angewandten Methode und z. B. Mathematica ist man rasch am Ziel. Man findet:

• *Calculer directement μ et σ^2 à la main est un peu fatigant. Par la méthode appliquée lors de la distribution binomiale et p.ex. Mathematica on trouve vite le résultat désiré. On trouve:*

Satz: • **Théorème:** $\mu = r \cdot \frac{1-p}{p}, \sigma^2 = r \cdot \frac{1-p}{p^2} = \frac{\mu}{p}$

Bemerkung: • **Remarque:** $\mu < \sigma^2 \leadsto$ negative Binomialverteilung • *distribution binomiale négative*
 $\mu > \sigma^2 \leadsto$ Binomialverteilung • *distribution binomiale*
 $\mu = \sigma^2 \leadsto$ Poissonverteilung • *distribution de Poisson*
 Für $r = 1$ erhält man die geometrische Verteilung. • *Pour $r = 1$ on obtient la distribution géométrique.*

Konsequenz: • **Conséquence:** Durch Schätzen von μ und σ^2 bei einer Stichprobe erhält man damit einen Hinweis auf den Typ der Verteilung.

• *L'estimation de μ et σ^2 d'un échantillon, on obtient une indication pour déterminer le type de la distribution.*

Zur Berechnung der Wahrscheinlichkeiten kann man sofort eine Rekursionsformel ablesen:

• *Pour calculer les probabilités, on peut tout de suite déduire une formule de récurrence:*

Formel: • **Formule:** $p_0 = p^r, p_{k+1} = \frac{r+k}{k+1} (1-p) \cdot p_k$

Bsp.: • **Exemple:**

Gegeben: Maschine, die durchschnittlich alle 4 Wochen (20 Arbeitstage) einmal eingesetzt wird. Eine Revision nach einem Einsatz dauert 3 Tage. Wahrscheinlichkeit, dass sie genau am 1. Tag ($r = 1, k = 0$) oder genau am 2. Tag ($r = 1, k = 1$) oder genau am 3. Tag ($r = 1, k = 2$) wieder gebraucht wird?

• *Donnée: Une machine qui est utilisée en moyenne tous les 4 semaines (20 journées de travail). Après une utilisation, on fait une révision du moteur qui dure 3 jours. Probabilité qu'elle sera utilisée exactement le 1. jour ($r = 1, k = 0$) ou exactement au 2. jour ($r = 1, k = 1$) ou exactement au 3. jour ($r = 1, k = 2$) pendant la révision?*

$$\leadsto p = \frac{1}{20}, q = 1 - p, P(X = k) = \binom{r+k-1}{r-1} \cdot p^r \cdot (1-p)^k, r = 1$$

$$\leadsto P(X = 0) = 0.05, P(X = 1) = 0.0475, P(X = 2) = 0.0428688.$$

4.8.5 Geometrische Verteilung — Distribution géométrique

Definition: • **Définition:** Die **geometrische Verteilung** ist der Spezialfall der Pascalverteilung für $r = 1$.

• *La **distribution géométrique** est le cas spécial de la distribution de Pascal pour $r = 1$.*

Es gilt daher: • *Il vaut donc:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Geometrische Verteilung • *Distribution géométrique*

Beh.: • Thè.:

$$p_k = P(X = k) = p \cdot (1 - p)^k = (1 - q) \cdot q^k, \quad k \in \mathbb{N}, \quad 0 < p < 1$$

$p = 1 - q$ ist die Wahrscheinlichkeit eines Misserfolgs, $q = 1 - p$ die Wahrscheinlichkeit eines Erfolgs und X gibt die zufällige Anzahl Erfolge bis zum 1. Misserfolg bei unabhängigen Versuchen. k sei der Parameter der Erfolge bis zum 1-ten **Misserfolg** ($k + 1$).

• $p = 1 - q$ est la probabilité d'un échec, $q = 1 - p$ la probabilité d'un succès et X donne le nombre aléatoire de succès jusqu'au premier échec. Les expériences partielles sont indépendantes. k soit le paramètre des succès jusqu'au 1-er **échec** ($k + 1$).

$$P(X \leq k) = P(X = 0) + \dots + P(X = k) = \sum_{j=0}^k (1 - q) \cdot q^j = (1 - q) \cdot \sum_{j=0}^k q^j = (1 - q) \cdot \frac{1 - q^{k+1}}{1 - q} = 1 - q^{k+1}$$

Satz: • Théorème: $P(X \leq k) = 1 - q^{k+1} \rightsquigarrow$ Geometrische Folge! • *Suite géométrique!*

Aus der Pascalverteilung erhalten wir für $r = 1$ die Momente: • *On obtient les moments à l'aide de la distribution de Pascal pour $r = 1$:*

Satz: • Théorème: $\mu = \frac{1 - p}{p}, \quad \sigma^2 = \frac{1 - p}{p^2} = \frac{\mu}{p}$

Bemerkung: • Remarque:

Die geometrische Verteilung (wegen der geometrischen Folge) ist das diskrete Analogon zur Exponentialverteilung.

• *La distribution géométrique (à cause de la suite géométrique) prend la place de la distribution exponentielle aux distributions discrètes.*

Bsp.: • Exemple:

Gegeben: Maschine, die durchschnittlich alle 4 Wochen (20 Arbeitstage) einmal eingesetzt wird. Eben ist sie eingesetzt worden. Wahrscheinlichkeit, dass sie innerhalb der nächsten 10 Tagen gebraucht wird?

• *Donnée: Une machine qui est utilisée en moyenne tous les 4 semaines (20 journées de travail). Maintenant on vient de l'utiliser. Probabilité qu'elle sera utilisée dans les prochains 10 jours?*

$$p = \frac{1}{20} = 0.05, \quad \mu = 20 - 1 = \frac{q}{1 - q} \Rightarrow q = \frac{19}{20} = 0.95. \quad (p = 1 - q = \frac{1}{20} = 0.05).$$

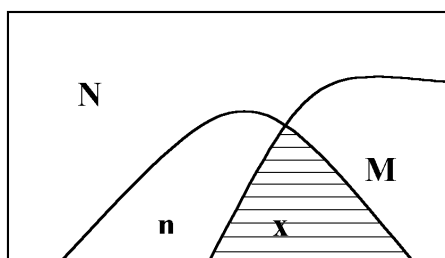
$$P(X \leq 9) = 1 - P(X \geq 10) = 1 - \left(\frac{1}{20}\right)^{9+1} = 1 - 0.95^{10} \approx 0.401263.$$

4.8.6 Hypergeometrische Verteilung — Distribution hypergéométrique

Im Unterschied zur Bernoulliverteilung sind bei der hypergeometrischen Verteilung die Experimente nicht unabhängig. Im Modell der Bernoulliverteilung ziehen wir mit zurücklegen, im Modell der hypergeometrischen Verteilung hingegen ziehen wir **ohne zurücklegen**. p ändert daher hier von Versuch zu Versuch.

• *Par contre à la distribution de Bernoulli, à la distribution hypergéométrique les expériences ne sont pas indépendantes. Dans le modèle de la distribution de Bernoulli, nous tirons un élément qui tout de suite après est remis à sa place originale. Dans le modèle de la distribution hypergéométrique par contre nous tirons **sans remettre l'élément**. Par conséquent p change ici d'une expérience à l'autre.*

Modellversuch: • Expérience modèle:



Modell: • *Modèle:*

Geg.: Gefäß mit N Elementen.

• *Donné: Récipient qui contient N éléments. M Elemente sind defekt.*

• *M éléments sont défectueux.*

n Elemente werden gezogen ohne zurückzulegen.

• *On tire n éléments sans les remettre.*

x der gezogenen n Elemente sind defekt.

• *x parmi les n éléments sont défectueux.*

Problem: • **Problème:**

Wahrscheinlichkeit, dass bei n mal ziehen x defekte Elemente kommen? • *Probabilité d'obtenir x éléments défectueux si on tire n éléments?*

g = günstige Fälle: Wähle x aus M und gleichzeitig $n - x$ aus $N - M$. • *g = cas favorables: Choisir x parmi M et au même temps $n - x$ parmi $N - M$.* $\leadsto g = \binom{M}{x} \cdot \binom{N - M}{n - x}$

m = mögliche Fälle: n aus N auswählen. • *m = cas possibles: Choisir n parmi N .* $\leadsto m = \binom{N}{n}$

$$\Rightarrow \frac{g}{m} = \frac{\binom{M}{x} \cdot \binom{N - M}{n - x}}{\binom{N}{n}}$$

Definition: • **Définition:**

Beim beschriebenen Modellversuch ist die **hypergeometrische Verteilung** gegeben durch:

• *A l'expérience modèle décrit, la **distribution hypergéométrique** est donnée par:*

$$f(x) = \frac{\binom{M}{x} \cdot \binom{N - M}{n - x}}{\binom{N}{n}}$$

Bemerkung: • **Remarque:**

Der Name „hypergeometrische Verteilung“ kommt von der sogenannten „hypergeometrischen Funktion“. • *La notion „distribution hypergéométrique“ vient de la „fonction hypergéométrique“.*

Wie früher schon kann man auch hier $G(t)$ verwenden, um μ und σ^2 zu berechnen. Die Rechnung ergibt:

• *Comme on l'a déjà exécuté dans d'autres cas, on peut utiliser aussi dans ce cas $G(t)$ pour calculer μ et σ^2 . Le calcul livre:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**Hypergeometrische Verteilung
• *Distribution hypergéométrique***Beh.:** • **Thè.:**

$$\mu = n \cdot \frac{M}{N}$$

$$\sigma^2 = \frac{n \cdot M \cdot (N - M) \cdot (N - n)}{N^2 \cdot (N - 1)}$$

Bemerkung: • **Remarque:**

1. Sind $N, M, N - M$ gross und n klein, so ist die hypergeometrische Verteilung eine Approximation der Binomialverteilung. $\leadsto p = \frac{M}{N}$
 • *Si $N, M, N - M$ sont grands et n est petit, la distribution hypergéométrique est une approximation de la distribution binomiale.* $\leadsto p = \frac{M}{N}$
2. Ist p klein, n gross und N gross gegenüber n , so ist die hypergeometrische Verteilung eine Approximation der Poissonverteilung. $\leadsto \mu = n \cdot p$
 • *Si p est petit, n grand et N grand par rapport à n , la distribution hypergéométrique est une approximation de la distribution de Poisson.* $\leadsto \mu = n \cdot p$

Bsp.: • **Exemple:**

Gegeben: Kiste mit 500 Kugellagern. Prüfung nach Vertrag: 50 (30) Lager entnehmen, untersuchen. Falls alle gut sind, wird die Lieferung akzeptiert. Sonst eben nicht. Nach Liefervertrag: Voraussichtlich maximal 2% $\hat{=}$ 10 sind Ausschuss. Berechne: $P(k = 0)$.

• *Donné: Une caisse qui contient 500 roulements à billes. L'examen selon le contrat: Retirer 50 (30) roulements, examiner. Si tous sont en ordre, la livraison est acceptée. Sinon, justement on ne les accepte pas. D'après le contract de livraison: Probablement maximal 2% $\hat{=}$ 10 sont de mauvaise qualité. Calculer: $P(k = 0)$.*

$$N = 500, \quad k = 0, \quad p = \frac{2}{100} = 0.02 = \frac{M}{N} = \frac{M}{500} \Rightarrow M = 500 \cdot 0.02 = 10.$$

$$n = 50 \Rightarrow P(X = k = 0) = \frac{\binom{M}{k} \cdot \binom{N-M}{n-k}}{\binom{N}{n}} = \frac{\binom{10}{0} \cdot \binom{500-10}{50-0}}{\binom{500}{50}} = 0.345162.$$

$$n = 30 \Rightarrow P(X = k = 0) = \frac{\binom{M}{k} \cdot \binom{N-M}{n-k}}{\binom{N}{n}} = \frac{\binom{10}{0} \cdot \binom{500-10}{30-0}}{\binom{500}{30}} = 0.535489.$$

4.8.7 Beispiel — Exemple

Bsp.: • **Exemple:**

Experiment: Eingespanntes Gewehr, 10 Schuss auf eine Scheibe, Trefferwahrscheinlichkeit pro Schuss = 0.1. Wie gross ist die Wahrscheinlichkeit, mindestens einen Treffer zu erzielen?

• **Expérience:** Fusil abloqué, 10 coups sur une cible, probabilité par coup de toucher le but = 0.1. Quelle est la probabilité de toucher au moins une fois?

Trick: • **Truc:** A : mindestens ein Treffer $\Rightarrow \bar{A}$: kein Treffer. • A : *toucher au moins une fois* $\Rightarrow \bar{A}$: *ne pas toucher*.

Sei • *Soit* A_0 : Treffer bei einem Schuss. • A_0 : *Toucher avec un seul coup*.

$$\leadsto P(A_0) = p = \frac{1}{10}, \quad P(A) = 1 - P(\bar{A}), \quad P(\bar{A}) = P(\bar{A}_0)^{10} \Rightarrow P(A) = 1 - P\left(\frac{9}{10}\right)^{10} \approx 0.651$$

$\leadsto$ Verteilung? • *Distribution?*

4.9 Spezielle stetige Verteilungen — Distributions continues spéciales

4.9.1 Allgemeines — Généralités

Sei $f(x)$ eine stückweise stetige Wahrscheinlichkeitsdichte mit $D_f = \mathbb{R}$ und $F(x)$ die zugehörige Verteilungsfunktion. Im folgenden wollen wir nun spezielle Varianten von $f(x)$ diskutieren.

• $f(x)$ soit une densité de probabilité pièce par pièce incessante avec $D_f = \mathbb{R}$ et $F(x)$ la fonction de distribution affiliée. Nous allons maintenant discuter des variantes spéciales de $f(x)$.

4.9.2 Rechtecksverteilung — Distribution rectangulaire

Sei $f(x)$ die Wahrscheinlichkeitsdichte einer Zufallsvariablen X .

• *Soit* $f(x)$ la fonction de densité d'une variable aléatoire X .

Die zu X gehörige Verteilung heisst **Rechtecksverteilung** oder **stetige gleichmässige Verteilung**, wenn gilt:

• *La distribution donnée par* X *s'appelle* **distribution rectangulaire**, *s'il vaut:*

Definition: • **Définition:**

$$f(x) = \begin{cases} \frac{1}{b-a} & a \leq x \leq b \\ 0 & x < a \vee x > b \end{cases}$$

Man rechnet: • *On calcule:*

$$1. \mu = m_1 = E(X^1) = \int_{-\infty}^{\infty} x^1 \cdot f(x) dx = \int_a^b x \cdot \frac{1}{b-a} dx = \frac{a+b}{2}$$

$$2. \sigma^2 = m_2 - m_1^2 = \int_{-\infty}^{\infty} x^2 \cdot f(x) dx - \mu^2 = \int_a^b x^2 \cdot \frac{1}{b-a} dx - \left(\frac{a+b}{2}\right)^2 = \frac{a^2 + ab + b^2}{3} - \frac{a^2 + 2ab + b^2}{4} = \frac{(b-a)^2}{12}$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Rechtecksverteilung

• *Distribution rectangulaire*

Beh.: • **Thè.:**

$$\mu = \frac{a+b}{2}, \quad \sigma^2 = \frac{(b-a)^2}{12}$$

Bsp.: • **Exemple:** Siehe Seite 92. • *Voir page 92.*

4.9.3 Normal- oder Gaussverteilung — Distribution normale ou de Gauss

Vorbemerkungen — Remarques préliminaires

Historisch gesehen ist die Normalverteilung schon von Gauss im Zusammenhang mit der Theorie der zufälligen (nicht systematische) Messfehler (Fehlerkurve) benutzt worden. Sie ist aus folgenden Gründen sehr wichtig:

• *Historiquement la distribution normale a déjà été utilisée par Gauss dans le contexte de la théorie des erreurs de mesure aléatoires, non-systématiques ou non-méthodiques (courbe des erreurs). Les raisons suivantes montrent son importance:*

1. In der Praxis sind viele Zufallsvariablen annähernd normalverteilt. Ein Problem bleibt aber die Tatsache, dass in der Praxis beliebig grosse Messgrössen wegen der physikalischen Beschränktheit nicht existieren.
• *Dans la pratique, beaucoup de variables aléatoires sont à peu près distribuées de façon normale. Reste le problème que, dans la pratique, les grandeurs infinies mesurées n'existent pas pour des raisons physiques.*
2. Viele Messgrössen lassen sich einfach auf die Normalverteilung transformieren oder sind durch solche approximierbar.
• *Beaucoup de grandeurs mesurées peuvent être transformées de façon simple en une distribution normale ou elles peuvent être approximées par une telle distribution.*
3. Die Normalverteilung spielt auch bei statistischen Prüfverfahren eine Rolle.
• *La distribution normale joue aussi un rôle aux méthodes de contrôle statistiques.*
4. Die Verteilung additiv überlagerter, zufälliger und unabhängiger Gesamtfehler hat als Grenzverteilung die Normalverteilung (**zentraler Grenzwertsatz**).
• *La distribution totale de la superposition d'erreurs aléatoires et indépendantes a comme distribution limite la distribution normale **théorème limite central**.*

Verteilungsfunktion — Fonction de distribution

Aus der Analysis wissen wir: • *De l'analyse on sait:* $\int_{-\infty}^{\infty} e^{-t^2} dt = \sqrt{\pi}$

Hinweis zum Beweis: • *Indication quant à la preuve:*

$$\text{Sei } \bullet \text{ Soit } I(R) = \int_{-\infty}^R e^{-t^2} dt \Rightarrow (I(R))^2 = \int_{-\infty}^R e^{-x^2} dx \cdot \int_{-\infty}^R e^{-y^2} dy = \int_{-\infty}^R \left(\int_{-\infty}^R e^{-(x^2+y^2)} dx \right) dy$$

Statt über das Quadrat $[0, R] \times [0, R]$ zu integrieren, kann man auch über den innern und äusseren Viertelskreis in Polarkoordinaten integrieren, die das Quadrat natürlich umschliessen. Für $R \rightarrow \infty$ erhält man dann den Grenzwert $\sqrt{\pi}$.

• Au lieu d'intégrer sur le carré $[0, R] \times [0, R]$, on peut aussi intégrer en coordonnées polaires sur les cercles intérieurs et extérieurs qui incluent naturellement le carré. Pour $R \rightarrow \infty$, on obtient donc la valeur limite $\sqrt{\pi}$.

Damit können wir folgern: • Nous pouvons en déduire:

$$\begin{aligned} \int_{-\infty}^{\infty} e^{-t^2} dt &= \sqrt{\pi} \Rightarrow \int_{-\infty}^{\infty} e^{-\frac{t^2}{2}} dt = \sqrt{2\pi} \Rightarrow \int_{-\infty}^{\infty} e^{-\frac{t^2}{2 \cdot s^2}} dt = s\sqrt{2\pi} \Rightarrow \int_{-\infty}^{\infty} e^{-\frac{(t-m)^2}{2 \cdot s^2}} dt = s\sqrt{2\pi} \\ \Rightarrow \frac{1}{s\sqrt{2\pi}} \cdot \int_{-\infty}^{\infty} e^{-\frac{(t-m)^2}{2 \cdot s^2}} dt &= 1, \quad D_f = \mathbb{R} \end{aligned}$$

Folgerung: • Conclusion:

$\frac{1}{s\sqrt{2\pi}} \cdot \int_{-\infty}^{\infty} e^{-\frac{(t-m)^2}{2 \cdot s^2}} dt = 1 \Rightarrow f(x) := \frac{1}{s\sqrt{2\pi}} \cdot e^{-\frac{(t-m)^2}{2 \cdot s^2}}$ ist als stetige Verteilungsfunktion zulässig.

• $\frac{1}{s\sqrt{2\pi}} \cdot \int_{-\infty}^{\infty} e^{-\frac{(t-m)^2}{2 \cdot s^2}} dt = 1 \Rightarrow f(x) := \frac{1}{s\sqrt{2\pi}} \cdot e^{-\frac{(t-m)^2}{2 \cdot s^2}}$ est admissible comme fonction de distribution continue.

Definition: • Définition:

Sei • Soit $f(x) := \varphi(x; m, s^2) = \frac{1}{s\sqrt{2\pi}} e^{-\frac{(x-m)^2}{2 \cdot s^2}}$

Sei • Soit $\varphi(x) := \varphi(x; 0, 1) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$

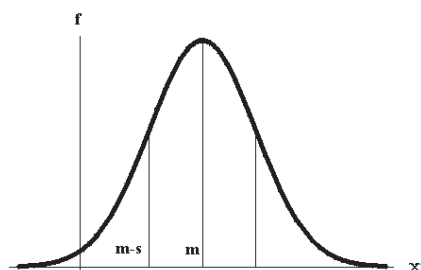
Die Dichtefunktion $f(x)$ definiert eine **Normalverteilung** oder **Gaussverteilung** mit den Parametern m und s zur Zufallsvariablen X .

• La fonction de densité $f(x)$ définit une **distribution normale** ou **distribution de Gauss** avec les paramètres m et s pour la variable aléatoire X .

Verteilungsfunktion: • Fonction de distribution:

$$F(x) := \Phi(x; m, s^2) = \frac{1}{s\sqrt{2\pi}} \cdot \int_{-\infty}^x e^{-\frac{(t-m)^2}{2 \cdot s^2}} dt$$

Symbol: • Symbole: $X \in N(m; s^2) \quad (\leadsto \quad Z = \frac{-(X-m)^2}{2 \cdot s^2} \in N(0; 1))$



Der Graph heisst auch **Gaussglocke**. Er ist symmetrisch zu $x = m$

• On appelle le graphe aussi **cloche de Gauss**. Il est symétrique à $x = m$

Für den Erwartungswert gilt: • Pour la valeur moyenne il vaut:

$$\mu = E(X) := \int_{-\infty}^{\infty} f(x) \cdot x dx = \int_{-\infty}^{\infty} x \cdot \frac{1}{s\sqrt{2\pi}} e^{-\frac{(t-m)^2}{2 \cdot s^2}} dx = m$$

Berechnung der Varianz (z.B. Einsatz von Mathematica): • *Calculer la variance (p.ex. employer Mathematica):*

$$\sigma^2 = E(X^2) - \mu^2 = G_t''(0) - \mu^2 = \frac{d^2}{dt^2} \left(\int_{-\infty}^{\infty} e^{xt} f(x) dx \right) \Big|_{t=0} - \mu^2 = s^2$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$X \in N(m; s^2)$$

Beh.: • **Thè.:**

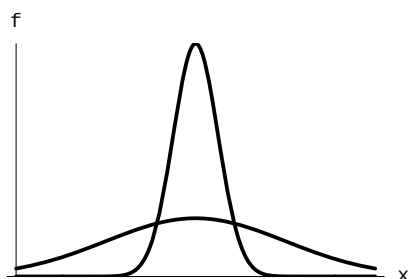
$$1. \mu = m$$

$$2. \sigma^2 = s^2$$

Definition: • **Définition:**

Ist $\mu = 0$ und $\sigma^2 = 1$, so hat man die **standardisierte Normalverteilung**.

• $S'\mu = 0$ et $\sigma^2 = 1$, on a la **distribution normale standardisée**.



Flache Kurve: $\sigma = 1$, steile Kurve: $\sigma = 25$.

• *Courbe plate: $\sigma = 1$, courbe raide: $\sigma = 25$.*

Bemerkung: • **Remarque:**

$$F(x) := \Phi(x; m, s^2) = \frac{1}{s\sqrt{2\pi}} \cdot \int_{-\infty}^x e^{\frac{-(t-m)^2}{2s^2}} dt$$

$$P(a < x \leq b) = F(b) - F(a)$$

$$\Phi(x; 0, 1) = \frac{1}{\sqrt{2\pi}} \cdot \int_{-\infty}^x e^{\frac{-(t)^2}{2}} dt$$

$$F(x) = \Phi(x; \mu, \sigma^2) = \Phi\left(\frac{x-\mu}{\sigma}; 0, 1\right) = \Phi(z; 0, 1)$$

$$\leadsto X \in N(\mu; \sigma^2) \Leftrightarrow Z = \frac{X - \mu}{\sigma} \in N(0; 1)$$

Die Verteilungsfunktion $F(x)$ ist nicht elementar berechenbar. Man verwendet dazu Rechner (Error function, *Erf*) oder Tabellen. Oft ist $\Phi(z; 0, 1)$ tabelliert. Damit lässt sich auch $\Phi(x; \mu, \sigma^2)$ bestimmen.

• *On ne peut pas calculer la fonction de répartition $F(x)$ de façon élémentaire. Pour arriver quand-même à un résultat, on utilise des calculateurs (error function, *Erf*) ou des tableaux. Souvent on trouve $\Phi(z; 0, 1)$ dans des tableaux. A l'aide de ces tableaux on peut trouver $\Phi(x; \mu, \sigma^2)$.*

Zwischen $\mu - \sigma$ und $\mu + \sigma$ liegt ca. 68% der Fläche unter der Kurve von $f(x)$, zwischen $\mu - 2\sigma$ und $\mu + 2\sigma$ liegt ca. 95.5% der Fläche und zwischen $\mu - 3\sigma$ und $\mu + 3\sigma$ liegt ca. 99.7% der Fläche.

• *Entre $\mu - \sigma$ et $\mu + \sigma$ on trouve environ 68% de la surface sous la courbe de $f(x)$, entre $\mu - 2\sigma$ et $\mu + 2\sigma$ on trouve environ 95.5% de la surface et entre $\mu - 3\sigma$ et $\mu + 3\sigma$ on trouve environ 99.7% de la surface.*

Weitere Informationen siehe unter: • *D'autres informations voir:*

<http://de.wikipedia.org/wiki/Normalverteilung>

4.9.4 Grenzwertsätze von Moivre Laplace — Théorèmes limites de Moivre Laplace

Diese Sätze behandeln die Annäherung einer Bernoulliverteilung durch eine Normalverteilung.

• *Ces théorèmes traitent l'approximation de la répartition de Bernoulli par une distribution normale.*

4.9.5 Lokaler Grenzwertsatz — Théorème limite locale

Geg.: • **Donné:**

Bernoulliverteilung • *Répartition de Bernoulli*

$$f_n(x) = \binom{n}{x} p^x q^{n-x} \rightsquigarrow \mu = n \cdot p, \sigma^2 = n \cdot p \cdot q, q = 1 - p, 0 < p < 1$$

Normalverteilung • *Distribution normale*

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma^2} e^{-\frac{z^2}{2}}, \mu = n \cdot p, \sigma^2 = n \cdot p \cdot q, q = 1 - p, z = \frac{x - \mu}{\sigma} \in \mathbb{R}$$

Wir wollen zeigen, dass für „kleine“ x $f_n(x)$ gleichmässig gegen $f(x)$ konvergiert für alle x in einem beliebigen endlichen Intervall. • *Nous voulons prouver, que pour des x "petits" $f_n(x)$ converge de façon uniforme vers $f(x)$ pour tous les x dans un intervalle quelconque et fini.*

Satz: • **Théorème:**

Vor.: • **Hyp.:**

Wie beschrieben • *Comme décrit*

Beh.: • **Thè.:**

$$f_n(x) \sim f(x)$$

Symbol: • **Symbole:** " $\sim$ ": lies **asymptotisch gleich** • " $\sim$ ": lire **égale de façon asymptotique**

Beweis: • **Preuve:**

Wir greifen auf die Formel von Stirling zurück (Seite 38):

• *Nous utilisons la formule de Stirling (page 38):*

$$\rightsquigarrow k! \approx \sqrt{2\pi k} \left(\frac{k}{e}\right)^k, k! = \sqrt{2\pi k} \cdot \left(\frac{k}{e}\right)^k \cdot e^{\frac{\Phi}{12k}}, 0 < \Phi < 1$$

Der Beweis dieser Formel ist mit unseren Mitteln ein grösseres Unterfangen, vgl. z.B. Amann u. Escher, Analysis II (Bibl.: A2) oder van der Waerden Bibl. A13.

• *La preuve de cette formule est une affaire étendue, voir p.ex. Amann u. Escher, Analysis II (Bibl.: A2) ou van der Waerden Bibl. A13.*

$$\begin{aligned} \rightsquigarrow f_n(x) &= \binom{n}{x} p^x q^{n-x} = \frac{n!}{x! \cdot (n-x)!} p^x q^{n-x} = \\ &= \sqrt{2\pi n} \cdot \left(\frac{n}{e}\right)^n \cdot e^{\frac{\Phi_n}{12n}} \cdot \frac{1}{\sqrt{2\pi x}} \cdot \left(\frac{x}{e}\right)^{-x} \cdot e^{-\frac{\Phi_x}{12x}} \cdot \frac{1}{\sqrt{2\pi(n-x)}} \cdot \left(\frac{n-x}{e}\right)^{-(n-x)} \cdot e^{-\frac{\Phi_{n-x}}{12(n-x)}} \cdot p^x \cdot q^{n-x} = \\ &= \frac{\sqrt{2\pi n}}{\sqrt{2\pi(n-x)} \cdot \sqrt{2\pi x}} \cdot \left(\frac{n}{e}\right)^n \cdot \left(\frac{e}{x}\right)^x \cdot \left(\frac{e}{n-x}\right)^{(n-x)} \cdot e^{\frac{\Phi_n}{12n}} \cdot e^{-\frac{\Phi_x}{12x}} \cdot e^{-\frac{\Phi_{n-x}}{12(n-x)}} \cdot p^x \cdot q^{n-x} = \end{aligned}$$

$$\frac{1}{\sqrt{2\pi}} \cdot n^{n+\frac{1}{2}} \cdot p^x \cdot \frac{1}{x^{x+\frac{1}{2}}} \cdot \frac{1}{(n-x)^{(n-x)+\frac{1}{2}}} \cdot q^{n-x} \cdot e^{\frac{\Phi_n}{12n} - \frac{\Phi_x}{12x} - \frac{\Phi_{(n-x)}}{12(n-x)}} =$$

$$\frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}} \cdot \left(\frac{n \cdot p}{x}\right)^{(x+\frac{1}{2})} \cdot \left(\frac{n \cdot q}{(n-x)}\right)^{(n-x+\frac{1}{2})} \cdot e^\tau = \frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}} \cdot e^{-h} \cdot e^\tau,$$

$$\tau = \frac{1}{12} \cdot \left(\frac{\Phi_n}{n} - \frac{\Phi_x}{x} - \frac{\Phi_{(n-x)}}{(n-x)}\right), \quad \Phi_i \in (0, 1), \quad h = \left(x + \frac{1}{2}\right) \cdot \ln\left(\frac{n \cdot p}{x}\right) + \left(n - x + \frac{1}{2}\right) \cdot \ln\left(\frac{n \cdot q}{(n-x)}\right)$$

Sei • Soit $z = \frac{x-\mu}{\sigma} = \frac{x-n \cdot p}{\sqrt{n \cdot p \cdot q}}$, $|z| < K = \text{const.}$, $p+q=1$

$$\leadsto h(x) := w(z) = \left(np + \frac{1}{2} + z\sqrt{npq}\right) \cdot \ln\left(1 + z\sqrt{\frac{q}{np}}\right) + \left(nq + \frac{1}{2} - z\sqrt{npq}\right) \cdot \ln\left(1 - z\sqrt{\frac{p}{nq}}\right)$$

n genügend gross • n assez grand $\leadsto u_1 = |z\sqrt{\frac{q}{np}}| < 1$, $u_2 = |z\sqrt{\frac{p}{nq}}| < 1$

$\leadsto \ln(1+u_1), \ln(1-u_2)$ sind daher in Potenzreihen um das Zentrum 1 entwickelbar:

• On peut développer $\ln(1+u_1), \ln(1-u_2)$ dans des séries de puissances autour du centre 1:

$$\ln(u) = u - \frac{u^2}{2} + \frac{u^3}{3} - \frac{u^4}{4} + \dots$$

$$\leadsto \ln\left(1 + z\sqrt{\frac{q}{np}}\right) = z\sqrt{\frac{q}{np}} - \frac{z^2}{2} \frac{q}{np} + \dots, \quad \ln\left(1 - z\sqrt{\frac{p}{nq}}\right) = -z\sqrt{\frac{p}{nq}} - \frac{z^2}{2} \frac{p}{nq} + \dots$$

$$\leadsto w(z) = \left(np + \frac{1}{2} + z\sqrt{npq}\right) \cdot \left(z\sqrt{\frac{q}{np}} - \frac{z^2}{2} \frac{q}{np} + \dots\right) + \left(nq + \frac{1}{2} - z\sqrt{npq}\right) \cdot \left(-z\sqrt{\frac{p}{nq}} - \frac{z^2}{2} \frac{p}{nq} + \dots\right)$$

Hier multiplizieren wir aus, ordnen nach Potenzen von z und vereinfachen ($q=1-p$, z.B. Mathematica verwenden), so erhalten wir:

• Ici, nous multiplions les terms, après nous les ordonnons d'après les puissances de z et nous simplifions ($q=1-p$, par exemple appliquer Mathematica). Ainsi nous obtenons:

$$w(z) = \left(\frac{q-p}{2\sqrt{npq}} \cdot z + z^2 \cdot \left(\frac{1}{2} - \frac{1}{2n(1-p)} - \frac{1}{4np(1-p)} - \frac{p}{2n(1-p)}\right) + \frac{1}{\sqrt{n}} \cdot R(n,p) \cdot z^3\right),$$

$$|R(n,p)| < \text{const.} (z \in (0,1))$$

$$\leadsto \lim_{n \rightarrow \infty} w(z) = \frac{1}{2} \cdot z^2, \quad z = \frac{x-\mu}{\sigma} = \frac{x-n \cdot p}{\sqrt{n \cdot p \cdot (1-p)}} \in [\alpha, \beta] \subset \mathbb{R}$$

$\leadsto$ gleichmässige Konvergenz • convergence uniforme

$$\leadsto e^{-h(x)} = e^{-w(z)} \rightarrow e^{z^2/2} \quad \text{für} \quad \bullet \quad \text{pour } n \rightarrow \infty$$

$$\alpha \leq z = \frac{x-n \cdot p}{\sqrt{n \cdot p \cdot (1-p)}} \leq \beta \Rightarrow x \geq np + \alpha \sqrt{npq} = np(1 + \alpha \sqrt{\frac{q}{np}}),$$

$$n-x \geq nq - \beta \sqrt{npq} = nq(1 - \beta \sqrt{\frac{p}{nq}})$$

$$\Rightarrow |\tau| = \left|\frac{1}{12} \cdot \left(\frac{\Phi_n}{n} - \frac{\Phi_x}{x} - \frac{\Phi_{(n-x)}}{(n-x)}\right)\right| < \frac{1}{12} \cdot \left(\frac{1}{n} + \frac{1}{x} + \frac{1}{(n-x)}\right)$$

$$< \frac{1}{12} \cdot \left(\frac{1}{n} + \frac{1}{np(1 + \alpha \sqrt{\frac{q}{np}})} + \frac{1}{nq(1 - \beta \sqrt{\frac{p}{nq}})}\right) \rightarrow 0 \quad \text{für} \quad \bullet \quad \text{pour } n \rightarrow \infty$$

$$\leadsto e^\tau \rightarrow 1 \quad \text{für} \quad \bullet \quad \text{pour } n \rightarrow \infty$$

$$\leadsto f_n(x) = \frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}} \cdot e^{-h} \cdot e^\tau \rightarrow \underbrace{\frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}}}_{=0} \cdot e^{z^2/2} \cdot \underbrace{e^0}_1 = 0, \quad f_n(x) \sim \frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}} \cdot e^{z^2/2}$$

Satz: • **Théorème:** **Lokaler Grenzwertsatz** • **Théorème limite locale**

Vor.: • **Hyp.:**

$$\alpha \leq z = \frac{x - n \cdot p}{\sqrt{n \cdot p \cdot (1 - p)}} \leq \beta, \quad (z \in [\alpha, \beta]),$$

$$n \rightarrow \infty$$

Beh.: • **Thè.:**

$$f_n(x) \sim \frac{1}{\sqrt{2\pi \cdot n \cdot p \cdot q}} \cdot e^{z^2/2}$$

4.9.6 Grenzwertsatz von De Moivre/ Laplace — Théorème limite de De Moivre/ Laplace

Der Grenzwertsatz von De Moivre/ Laplace ist ein Korollar des lokalen Grenzwertsatzes. Er ist ein Spezialfall des zentralen Grenzwertsatzes, der hier nicht behandelt wird (vgl. Lit.).

• *Le théorème de limite de De Moivre/ Laplace est un corollaire du théorème de limite locale. C'est un cas spécial du théorème limite centrale qui n'est pas traité ici (voir lit.).*

Sei [...] die Gaußklammerfunktion (Floor, Int). • *Soit [...] la fonction de parenthèses de Gauss (Floor, Int).*

Satz: • **Théorème:** **Grenzwertsatz von De Moivre/ Laplace**

• *Théorème de limite de De Moivre/ Laplace*

Vor.: • **Hyp.:**

$$n \rightarrow \infty,$$

$$\Phi(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x e^{-u^2/2} du, \quad a < b, \quad \alpha = \left[\frac{a - n \cdot p - 0.5}{\sqrt{n \cdot p \cdot (1 - p)}} \right], \quad \beta = \left[\frac{b - n \cdot p + 0.5}{\sqrt{n \cdot p \cdot (1 - p)}} \right],$$

$$\alpha \leq z = \frac{x - n \cdot p}{\sqrt{n \cdot p \cdot (1 - p)}} \leq \beta$$

Beh.: • **Thè.:**

$$P(a \leq X \leq b) = \sum_{x=a}^b \binom{n}{x} p^x q^{n-x} \sim \Phi(\beta) - \Phi(\alpha) = \frac{1}{\sqrt{2\pi}} \int_{\alpha}^{\beta} e^{-u^2/2} du$$

4.9.7 Das Gesetz von Bernoulli der grossen Zahlen — La loi des grands nombres de Bernoulli

Das **Gesetz der grossen Zahlen** liefert eine Wahrscheinlichkeitsaussage zur Vermutung, dass sich die relative Häufigkeit für grosse n verhält wie die Wahrscheinlichkeit.

• *La loi des grands nombres fournit une explication propositionnelle pour la supposition que la fréquence suit la probabilité pour les grand n .*

Vor.: • **Hyp.:**

Sei A : Ereignis mit der Wahrscheinlichkeit $p \in (0, 1)$ bei einem Zufallsexperiment.

Sei X = Anzahl des Eintreffens von A beim n -maligen Ausführen des Experiments.

Sei $\varepsilon \in \mathbb{R}^+$.

• *Soit A : événement avec la probabilité $p \in (0, 1)$ à une expérience de hasard.*

Soit X =: nombre de réalisations de A si on exécute l'expérience n fois.

Soit $\varepsilon \in \mathbb{R}^+$.

Satz: • **Théorème:** **Bernoulli**

Beh.: • **Thè.:**

$$P\left(\left|\frac{x}{n} - p\right| < \varepsilon\right) \xrightarrow{n \rightarrow \infty} 1$$

Bemerkung: • **Remarque:**

Der Satz besagt also: Die Wahrscheinlichkeit oder Chance, dass die Abweichung der relativen Häufigkeit von der Wahrscheinlichkeit $p = P(X)$ kleiner als jedes beliebige positive ε wird, wenn nur n genügend gross ist, konvergiert gegen 1.

Achtung: Nur die Wahrscheinlichkeit konvergiert, nicht aber etwa die Differenz $|\frac{x}{n} - p|$. Eine grosse Wahrscheinlichkeit bedeutet nicht etwa Sicherheit!

• *Le théorème affirme donc: La probabilité ou la chance que l'écart de la fréquence relative de la probabilité $p = P(X)$ soit plus petit que n'importe quel nombre ε positif quelconque si n est suffisamment grand, converge vers 1. Attention: Seulement la probabilité converge, mais ne pas la différence $|\frac{x}{n} - p|$. Une grande probabilité non signifie pas sécurité!*

Wir können die Behauptung aus dem Satz von De Moivre/ Laplace folgern:

• *Nous pouvons déduire l'assertion du théorème de De Moivre/ Laplace:*

Beweis: • **Preuve:**

$$\begin{aligned} \left|\frac{x}{n} - p\right| < \varepsilon &\Leftrightarrow -\varepsilon < \frac{x}{n} - p < \varepsilon \Leftrightarrow p - \varepsilon < \frac{x}{n} < p + \varepsilon \Leftrightarrow a := (p - \varepsilon) \cdot n < x < (p + \varepsilon) \cdot n := b \\ \Rightarrow P\left(\left|\frac{x}{n} - p\right| < \varepsilon\right) &= P((p - \varepsilon) \cdot n < x < (p + \varepsilon) \cdot n) = P(a < x < b) = P(a \leq x \leq b), \end{aligned}$$

$$\text{De Moivre/ Laplace} \rightsquigarrow P(a \leq x \leq b) = \frac{1}{\sqrt{2\pi}} \int_{\alpha}^{\beta} e^{-u^2/2} du \xrightarrow{n \rightarrow \infty} \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-u^2/2} du$$

Denn es gilt: • *Car il vaut:*

$$\begin{aligned} \alpha &= \left[\frac{a - n \cdot p - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{(p - \varepsilon) \cdot n - n \cdot p - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{-\varepsilon \cdot n - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{-\varepsilon \cdot \sqrt{n}}{\sqrt{p \cdot q}} + \frac{0.5}{\sqrt{p \cdot q}} \right] \rightarrow -\infty, \\ \text{beta} &= \left[\frac{b - n \cdot p - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{(p + \varepsilon) \cdot n - n \cdot p - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{+\varepsilon \cdot n - 0.5}{\sqrt{n \cdot p \cdot (1-p)}} \right] = \left[\frac{+\varepsilon \cdot \sqrt{n}}{\sqrt{p \cdot q}} + \frac{0.5}{\sqrt{p \cdot q}} \right] \rightarrow \infty \end{aligned}$$

4.9.8 Bemerkung zum Zufall — Remarques quant au hasard

Wie wir eingangs gesehen haben, wird beim absoluten Zufall ein Grund ausgeschlossen, beim relativen Zufall ist ein Grund jedoch bloss nicht wahrnehmbar, existiert jedoch. Nun macht es aber keinen Sinn über die Existenz einer nicht wahrnehmbaren Sache zu rätseln. Man kann da nur noch glauben, nicht glauben oder eben die Frage Frage sein lassen.

Wir können daher von folgendem Begriffsverständnis ausgehen: Für den beobachtenden Menschen ist es Zufall, wenn das Eintreten eines Ereignisses unvorhergesehen, unbeabsichtigt ist, ohne erkennbare Ursache oder Gesetzmässigkeit. Zufall ist es, wenn das Eintreten eines Ereignisses nicht notwendig

aus einem gegebenen Gesamtereignis folgt, wenn alles auch hätte anders verlaufen können. Bloss eine Wirkung ist sichtbar, eine Ursache aber nicht.

Bei grösseren Ereignismengen werden aber nun dennoch Gesetzmässigkeiten sichtbar. Die relativen Häufigkeiten von Experimenten nähern sich den aus Modellen herleitbaren theoretischen Wahrscheinlichkeiten an (Gesetz der grossen Zahlen). Zufällige Ursachen haben hier also gesetzmässige, empirisch beobachtbare Wirkungen. Man kann daher trotzdem die sinnvolle Frage stellen, ob die Resultate nicht doch aus einem komplexen Wirkungszusammenhang gesetzmässig folgen. Die Vorhersage des Resultates könnte daher nur praktisch zu kompliziert, theoretisch aber möglich sein. (Z.B. Galton–Brett–Experiment!). Zufall wäre daher eine Folge der Komplexität. Die Sache kann aber auch umgekehrt sein: Relativ exakt bekannte Ursachen haben manchmal trotz bekannter Naturgesetze eine recht zufällige Wirkung. Z.B. bei der Wettervorhersage oder bei der Berechnung der Erdposition im Sonnensystem vor Jahrtausenden. Die unvermeidlichen Messfehler können sich bei Berechnungen derart fortpflanzen, dass der Endfehler den Wertebereich des Resultats übersteigt. Kleinste Veränderungen der Ursachen können grösste Abweichungen in der Wirkung zur Folge haben (Schmetterlingseffekt).

Man denke in diesem Zusammenhang an die Ausschweifungen des Determinismus (Laplace'scher Dämon, Negation des freien Willens) oder an die Unschärferelation (gesetzmässige Beschreibung des Unbekannten).

• *Comme nous avons vu au début, au hasard absolu une raison est exclue, par contre au hasard relatif, une raison n'est cependant pas perceptible, mais elle existe. Or, cela n'a aucun sens de se poser des questions au sujet de l'existence d'une chose qui n'est pas perceptible. On peut croire, ne pas croire ou laisser la question ouverte.*

Par conséquent nous pouvons commencer avec la compréhension de l'idée suivante: Pour l'être humain qui observe, c'est un hasard, si la réalisation d'un événement est imprévue, pas voulue, sans raison reconnaissable ni régularité. C'est un hasard si l'événement ne résulte pas nécessairement d'un ensemble d'événement donnés et si tout aurait pu se dérouler différemment. Seulement l'effet est visible, mais non pas une cause.

Pourtant, aux ensembles d'événements plus grands, des lois deviennent maintenant visibles. Les fréquences relatives d'expériences s'approchent aux probabilités déduisibles de façon théorique de modèles (loi des grands nombres). Les observations empiriques et accidentelles ont ici des effets qui suivent des lois. Par conséquent on peut poser quand même la question raisonnable si les résultats ne résultent pas suivant des lois dans un contexte d'effets complexe. Par conséquent la prédiction du résultat pourrait être pratiquement trop compliquée, tandis qu'elle est théoriquement possible. (P.ex. expérience avec la planche de Galton). Le hasard serait donc une suite de la complexité. Mais la chose peut être aussi inversée: Des causes connues relativement exactement ont parfois un effet bien accidentel malgré la loi de la nature connue. P.ex. à la prévision météorologique ou au calcul de la position que la terre avait dans le système solaire il y a des millénaires. Les erreurs de mesure inévitables peuvent tellement s'amplifier durant les calculs que l'erreur finale surmonte le domaine de valeur du résultat. Par conséquent les changements les plus petits des causes (constellations originales) peuvent avoir des écarts énormes dans l'effet (effet de papillon).

Qu'on pense dans ce contexte aux débauches du déterminisme (démon de Laplace, négation de la volonté libre) ou au rapport du flou (description déterminée de l'inconnu).

4.9.9 Tschebyscheffsche Ungleichung — Inéquation de Tschebyscheff

Sei Y eine beliebige Zufallsgrösse, die nicht normalverteilt sein muss. $E(Y)$ sei der Erwartungswert, $Var(Y)$ sei die Varianz und ε sei eine beliebige positive Zahl. Dann gilt (ohne Beweis, vgl. Anhang):

• *Soit Y une variable aléatoire quelconque, qui ne doit pas être distribuée de façon normale. Soit $E(Y)$ la valeur d'espérance, $Var(Y)$ soit la variance et ε soit un nombre positif et quelconque. Alors il vaut (sans preuve, voir annexe):*

Formel: • **Formule:**

Tschebyscheff

$$P(|Y - E(Y)| \geq \varepsilon) \leq \frac{Var(Y)}{\varepsilon^2}$$

4.9.10 Logarithmische Normalverteilung — Distribution normale logarithmique

Sei • Soit $f(t) := \varphi(t; \mu_L, \sigma_L^2) = \frac{1}{\sigma_L \sqrt{2\pi}} e^{\frac{-(t-\mu_L)^2}{2\sigma_L^2}}$

= Dichtefunktion der Normalverteilung, Parameter μ_L und σ_L , Zufallsvariablen Y .

• = fonction de densité $f(y)$ de la distribution normale, paramètres μ_L et σ_L , variable aléatoire Y .

$\leadsto$ Verteilungsfunktion: • Fonction de distribution: $F(y) := \Phi(y; \mu_L, \sigma_L^2) = \frac{1}{\sigma_L \sqrt{2\pi}} \cdot \int_{-\infty}^y e^{\frac{-(t-\mu_L)^2}{2\sigma_L^2}} dt$

$\leadsto Y$ ist normalverteilt. • Y est distribué de façon normale.

Sei • Soit $X = \log(Y)$, $x = \log(y)$ (Substitution) • (Substitution)

Falls in $f(y)$ einfach $y = \log_a(x)$, $a^y = x$ gesetzt wird, entspricht das Vorgehen nicht der bekannten Substitutionsregel. Wegen der Normierung muss gelten:

• Si on remplace $y = \log_a(x)$, $a^y = x$ dans $f(y)$, ça ne correspond pas à la règle de la substitution. (On n'a pas tenu compte de la dérivée intérieure). À cause de la normalisation il doit valoir:

$$1 = \int_{-\infty}^{\infty} f(y) dy = \int_{x=a^{-\infty}}^{x=a^{\infty}} f(\log_a(x)) \frac{dy}{dx} dx = \int_0^{\infty} f(\log_a(x)) \frac{\log_a(e)}{x} dx$$

Definition: • **Définition:**

Gegeben sei die Dichtefunktion: • Soit donnée la fonction de densité:

$$h(x) = \begin{cases} 0 & x \leq 0 \\ \frac{\log_a(e)}{x} \varphi(\log_a(x); \mu_L, \sigma_L^2) & x > 0 \end{cases}$$

Die durch $h(x)$ gegebene Verteilung heisst **logarithmische Normalverteilung** oder **Logonormalverteilung**.

• La distribution donnée par $h(x)$ s'appelle **distribution normale logarithmique**.

Die Verteilungsfunktion erhält man aus der Substitution von $F(y) := \Phi(y; \mu_L, \sigma_L^2)$: • Nous obtenons la fonction de distribution par la substitution de $F(y) := \Phi(y; \mu_L, \sigma_L^2)$:

Formel: • **Formule:**

$$H(x) = \frac{1}{\sigma_L \sqrt{2\pi}} \cdot \int_0^x e^{\frac{-(\log_a(z)-\mu_L)^2}{2\sigma_L^2}} \frac{\log_a(e)}{t} dt$$

Für $a = e$ erhält man durch Berechnung von Erwartungswert, Streuung und Median für die Logonormalverteilung:

• Pour $a = e$ on obtient par calcul de la valeur d'espérance, de la variance et du médian pour la distribution normale logarithmique:

Formel: • **Formule:**

$$\begin{aligned} E(X) &= e^{\mu_L + \sigma_L^2/2} \\ \text{Var}(X) &= e^{2\mu_L + \sigma_L^2} (e^{\sigma_L^2} - 1) \\ x_{0.5} &= e^{\mu_L} \end{aligned}$$

Anwendung: Zeitstudien, Lebensdaueranalysen. • Application: études des temps, analyse de longévité

Normalverteilung: Additive Überlagerungen $\leadsto$ Logonormalverteilung: Multiplikative Überlagerungen.

- *Distribution normale: Superpositions additives $\rightsquigarrow$ distribution normale logarithmique: superpositions multiplicatives.*

4.9.11 Exponentialverteilung — Distribution exponentielle

Es gilt: • *Il vaut:*

$$\int_0^{\infty} \alpha e^{-\alpha x} dy = -\frac{\alpha}{\alpha} e^{-\alpha x} \Big|_0^{x=\infty} = -(e^{-(\infty)} - e^{-(0)}) = -(0 - 1) = 1$$

Definition: • **Définition:**

Die stetige Zufallsvariable X genügt einer **Exponentialverteilung** mit dem Parameter α , wenn sie folgende Dichtefunktion besitzt:

- *La variable aléatoire X satisfait une **distribution exponentielle** avec le paramètre α , si elle a la fonction de densité suivante:*

$$f(x) = \begin{cases} 0 & x < 0 \\ \alpha e^{-\alpha x} & x \geq 0 \end{cases}$$

$\rightsquigarrow$ Verteilungsfunktion: • *Fonction de distribution:*

$$\int_{-\infty}^x f(t) dt = \int_0^x \alpha e^{-\alpha t} dt = 1 - e^{-\alpha x}$$

Formel: • **Formule:**

$$F(x) = \int_0^x f(t) dt = \begin{cases} 0 & x < 0 \\ 1 - e^{-\alpha x} & x \geq 0 \end{cases}$$

Erwartungswert, Streuung und Median: • *Valeur d'espérance, variance et médian:*

$$\mu = E(X) = \int_0^{\infty} x \cdot \alpha e^{-\alpha x} dx = x \cdot \frac{\alpha}{-\alpha} e^{-\alpha x} \Big|_0^{\infty} - \int_0^{\infty} (-1) e^{-\alpha x} dx = 0 - 0 + \int_0^{\infty} e^{-\alpha x} dx = \frac{1}{-\alpha} e^{-\alpha x} \Big|_0^{\infty} = \frac{1}{\alpha}$$

$$\sigma = Var(X) = \int_0^{\infty} x^2 \cdot \alpha e^{-\alpha x} dx - \mu^2 = \dots = \frac{2}{\alpha^2} - \frac{1}{\alpha^2} = \frac{1}{\alpha^2}$$

$$F(x_{0.5}) = 1 - e^{-\alpha x_{0.5}} = \frac{1}{2} \Rightarrow \dots \Rightarrow x_{0.5} = \frac{1}{\alpha} \ln(2)$$

Anwendung: Zeitmessungen (Atomzerfall, Arbeitszeiten, ...), Lebensdauer.

- *Application: mesure du temps (désintégration de l'atome, heures de travail, ...), durée de vie.*

4.9.12 Weibullverteilung — Distribution Weibull

Definition: • **Définition:** Eine stetige Zufallsvariable X besitzt eine dreiparametrische **Weibull-Verteilung** mit den Parametern $a > 0$, $b > 0$ und c , wenn ihre Dichtefunktion wie folgt gegeben ist:

• *Une variable aléatoire continue X a une **distribution de Weibull** aux trois paramètres $a > 0$, $b > 0$ et c , si la fonction de densité est donnée comme il suit:*

$$f(x) = \begin{cases} 0 & x \leq c \\ \frac{b}{a} \left(\frac{x-c}{a}\right)^{b-1} e^{-\left(\frac{x-c}{a}\right)^b} & x > c \end{cases}$$

$\leadsto$ Verteilungsfunktion: • *Fonction de distribution:*

Formel: • **Formule:**

$$F(x) = \begin{cases} 0 & x \leq c \\ 1 - e^{-\left(\frac{x-c}{a}\right)^b} & x > c \end{cases}$$

$c = 0 \leadsto$ Zweiparametrische Verteilung • *Distribution à deux paramètres*

$a = 1 \leadsto$ Reduzierte Verteilung • *Distribution réduite*

Bekannt aus der Analysis: • *Connue de l'analyse:*

Symbol: • **Symbole:** Sei $\Gamma(x)$ die **Gamma-Funktion**. • *Soit $\Gamma(x)$ la fonction Gamma.*

Definition: • **Définition:** $\Gamma(x) := \int_0^{\infty} e^{-t} t^{x-1} dt, \quad x > 0$

Formel: • **Formule:** $\Gamma(1) = 1, \quad \Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}, \quad \Gamma(x) = (x-1)\Gamma(x-1)$

Durch Berechnung erhält man: • *On obtient en calculant:*

Formel: • **Formule:**

$$\mu = E(X) = c + a \cdot \Gamma\left(1 + \frac{1}{b}\right)$$

$$\sigma^2 = Var(X) = a^2 \left(\Gamma\left(1 + \frac{2}{b}\right) - \Gamma^2\left(1 + \frac{1}{b}\right)\right)$$

$$x_{0.5} = c + a (\ln(2))^{1/b}$$

Anwendung: Lebensdauer- und Zuverlässigkeitsanalysen. • *Application: analyses de durée de vie et de fiabilité.*

4.9.13 Gammaverteilung — Distribution gamma

Definition: • **Définition:**

Eine stetige Zufallsvariable X besitzt eine **Gammaverteilung** mit den Parametern $b > 0$ und $p > 0$, wenn ihre Dichtefunktion wie folgt gegeben ist:

• *Une variable aléatoire continue X a une **distribution gamma** aux paramètres $a > 0$, $b > 0$ et c , si la fonction de densité est donnée comme il suit:*

$$f(x) = \begin{cases} 0 & x \leq 0 \\ \frac{b^p}{\Gamma(p)} x^{p-1} e^{-bx} & x > 0 \end{cases}$$

Bemerkung: • **Remarque:**

Die Exponentialverteilung ergibt sich als Spezialfall der Gammaverteilung für $b = \alpha$ und $p = 1$. Für $p \in \mathbb{N}$ heisst die Verteilung auch **Erlangverteilung**.

• *On obtient la distribution exponentielle comme cas spécial de la distribution gamma pour $b = \alpha$ et $p = 1$. Pour $p \in \mathbb{N}$ la distribution s'appelle aussi **distribution de Erlang**.*

Formel: • **Formule:**

$$\mu = E(X) = \frac{p}{b}, \quad \sigma^2 = \text{Var}(X) = \frac{p}{b^2}$$

Anwendung: Zuverlässigkeitsanalysen (Lebensdauerverteilung), Bedienungszeitenverteilung.

• *Application: analyses de fiabilité (distribution de durée de vie) distribution des temps de service.*

4.9.14 Ausblick — Autres distributions

Um weitere wichtige Verteilungen behandeln zu können, müssen wir erst Kenntnisse der mehrdimensionalen Zufallsgrößen (Zufallsvektoren) erwerben. Wichtig sind die Prüfverteilungen für statistische Tests.

• *Pour pouvoir traiter d'autres distributions importantes, nous devons acquérir des connaissances sur les variables aléatoires aux dimensions supérieures (vecteurs aléatoires). Les lois usuelles pour les tests sont importantes quant aux tests statistiques.*

4.10 Zufallsvektoren und deren Verteilung — Vecteurs aléatoires et leurs distributions

$\leadsto$ Mehrdimensionale Wahrscheinlichkeit • *Probabilité multidimensionnelle*

4.10.1 Fragestellung, Begriffe — Question, notions

Zufallsvektor, Verteilungsfunktion — Vecteur aléatoire, fonction de répartition

Bis jetzt haben wir **eindimensionale Zufallsgrößen** X betrachtet, die Werte $\in \mathbb{R}$ annehmen können. Die Verteilungsfunktion ist dann eine Funktion einer unabhängigen Variablen.

In einem Zufallsexperiment können jedoch verschiedene Ereignisse eintreten, die unabhängig oder abhängig sein können. Ebenso können zur Beschreibung eines Zufallsexperimentes mehrere Zufallsvariablen notwendig sein, die unabhängig oder abhängig sein können. Wir wollen hier den Modellfall von zwei Zufallsvariablen kurz studieren. ($N = 2$.)

• *Jusqu'à maintenant nous avons étudié les **variables aléatoires unidimensionnelles** X , qui*

atteignent des valeurs $\in \mathbb{R}$. La fonction de répartition est donc une fonction une variable indépendante. Par contre, dans une expérience aléatoire des événements différents peuvent être réalisés. Ces événements peuvent être indépendants ou dépendants. Également dans une expérience aléatoire plusieurs variables aléatoires peuvent être nécessaires à la description. Ces variables peuvent être indépendantes ou dépendantes. Ici nous allons étudier brièvement le cas modèle de deux variables aléatoires. ($N = 2$.)

Bsp.: • Exemple:

Bei der Qualitätskontrolle eines Loses von Zahnradwellen werden die Dicke X und gleichzeitig am selben Stück jeweils auch die Länge Y kontrolliert. Dicke und Länge werden jeweils in einem separaten Arbeitsgang gefertigt. Sie können daher als unabhängig angesehen werden. Da die Messungen jedoch paarweise vorgenommen werden, hängen Y und X dennoch irgendwie zusammen. Das führt zu Spezialfragen der mathematischen Statistik, auf die hier nicht eingegangen werden kann.

• *Au contrôle de qualité d'un lot d'arbres de roues dentées on contrôle l'épaisseur X et simultanément à la même pièce chaque fois aussi la longueur Y . L'épaisseur et la longueur sont fabriquées chaque fois pendant une phase de travail séparée. Par conséquent on peut les considérer comme indépendantes. Comme les mesurages cependant sont faits en paires, Y et X ont quand-même un rapport. Ça mène à des questions spéciales de la statistique mathématique que nous ne pouvons pas traiter ici.*

Seien X und Y zwei diskrete oder kontinuierliche Zufallsvariablen, die bei einem Zufallsexperiment auftreten. $\leadsto X$ kann Werte annehmen in $\{x_1, x_2 \dots x_k \dots\}$ oder in $\mathbb{R}$ (allgemein $x \in \mathbb{R}$) und Y Werte in $\{y_1, y_2 \dots y_k \dots\}$ oder in $\mathbb{R}$ (allgemein $y \in \mathbb{R}$).

• *Soyent X et Y deux variables aléatoires discrètes ou continues, qui apparaissent à une expérience de hasard. $\leadsto X$ peut prendre des valeurs dans $\{x_1, x_2 \dots x_k \dots\}$ ou dans $\mathbb{R}$ (généralement $x \in \mathbb{R}$) et Y peut prendre des valeurs dans $\{y_1, y_2 \dots y_k \dots\}$ ou dans $\mathbb{R}$ (généralement $y \in \mathbb{R}$).*

Definition: • Définition:

Der Vektor $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ heisst **zweidimensionaler Zufallsvektor**.

Entsprechend ist $\vec{X} = \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix}$ ein **n -dimensionaler Zufallsvektor**

• *Le vecteur $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ s'appelle vecteur aléatoire bidi-*

mensionnel. Correspondamment, $\vec{X} = \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix}$ s'appelle vecteur aléatoire de dimension n .

Sei $A = \{\omega\}$ ein Ereignis, das bei einem Zufallsexperiment mit zwei Zufallsvariablen eintritt. Dann nimmt X den Wert $X(A) = x$ an und gleichzeitig nimmt Y den Wert $Y(A) = y$ an. D.h. $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$

nimmt den Wert $\begin{pmatrix} y \\ x \end{pmatrix}$ (oder kurz das Wertepaar (x, y)) an. $\begin{pmatrix} y \\ x \end{pmatrix}$ ist dann eine **Realisierung** von $\vec{X}$.

• *Soit $A = \{\omega\}$ un événement, qui soit réalisé lors d'une expérience de hasard avec deux variables aléatoires. X y prend la valeur $X(A) = x$ et en même temps c'est Y qui prend la valeur $Y(A) = y$.*

Ça veut dire que $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ prend la valeur $\begin{pmatrix} y \\ x \end{pmatrix}$ (ou brièvement la paire (x, y)). $\begin{pmatrix} y \\ x \end{pmatrix}$ est donc une réalisation de $\vec{X}$.

Bemerkung: • **Remarque:**

Statt von einem Paar (X, Y) von Zufallsvariablen zu sprechen, kann man auch von einem **Wahrscheinlichkeitsvektor** $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ sprechen. Der Wahrscheinlichkeitszusammenhang des Paares (X, Y) (Wahrscheinlichkeitsfunktion) ist dann gegeben durch eine Doppelfolge $\langle (x_i, y_k, p_{i,k}) \rangle$.

• *Au lieu de considérer une paire (X, Y) de variables aléatoires, on peut parler aussi d'un vecteur aléatoire $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$. La loi de probabilité bi-dimensionnelle du couple (X, Y) (fonction aléatoire) est donc donnée par la double suite $\langle (x_i, y_k, p_{i,k}) \rangle$*

(Für die Entwicklung der Theorie beschränken wir uns hier auf $N = 2$. Die Übertragung auf $N > 2$ ist problemlos.)

• *(Pour le développement de la théorie nous nous limitons ici à $N = 2$. Le transfert à $N > 2$ est sans problèmes.)*

Für die zugehörigen Wahrscheinlichkeiten definieren wir:

• *Pour les probabilités affiliés nous définissons:*

Definition: • **Définition:**

$$F(x, y) = P(X \leq x \wedge Y \leq y), \quad x \in \mathbb{R}, \quad y \in \mathbb{R}$$

Im diskreten Fall ist speziell:

• *Dans le cas discret il est spécialement:*

$$p_{ik} := P(X = x_i \wedge Y = y_k)$$

$\leadsto F(x, y)$ ist eine Funktion $\mathbb{R}^2 \xrightarrow{F} [0, 1]$. • *$F(x, y)$ est une fonction $\mathbb{R}^2 \xrightarrow{F} [0, 1]$.*

Eigenschaften:• **Qualités:**

1. F ist in X und in Y monoton wachsend sowie rechtsseitig stetig.
 - *F est croissant de façon monotone dans X et dans Y et en plus continue depuis la droite.*
2. $F(-\infty, y) = 0, \quad F(x, -\infty) = 0, \quad F(\infty, \infty) = 1$

$$\text{Diskreter Fall: } \bullet \text{ Cas discret: } \sum_{i,k} p_{ik} = 1$$

3. $x_1 < x_2 \wedge y_1 < y_2 \Rightarrow F(x_2, y_2) - F(x_2, y_1) - F(x_1, y_2) + F(x_1, y_1) \leq 0$
4. $P(X \leq x) = P(X \leq x, Y \leq \infty) = F(x, \infty)$
 $P(Y \leq y) = P(X \leq \infty, Y \leq y) = F(\infty, y)$

Die ersten beiden Eigenschaften sind unmittelbar einsichtig. Bei der vierten Eigenschaft handelt es sich um eine **Spezialisierung** (Projektion oder Restriktion von $D_F = \mathbb{R}^N$ auf $\mathbb{R}^{N-1}$). Die dritte folgt wegen:

• *Les deux premiers propriétés sont évidentes. Quant à la quatrième propriété il s'agit d'une spécialisation (projection ou restriction de $D_F = \mathbb{R}^N$ à $\mathbb{R}^{N-1}$). La troisième propriété se voit à cause de:*

$$\begin{aligned}
& F(x_2, y_2) - F(x_2, y_1) - F(x_1, y_2) + F(x_1, y_1) \\
&= (P(X \leq x_2, Y \leq y_2) - P(X \leq x_2, Y \leq y_1)) - (P(X \leq x_1, Y \leq y_2) - P(X \leq x_1, Y \leq y_1)) \\
&= \underbrace{(P(X \leq x_2, Y \leq y_2) - P(X \leq x_2, Y \leq y_1))}_{=P(X \leq x_2, y_1 < Y \leq y_2)} - \underbrace{(P(X \leq x_1, Y \leq y_2) - P(X \leq x_1, Y \leq y_1))}_{=P(X \leq x_1, y_1 < Y \leq y_2)} \\
&= P(X \leq x_2, y_1 < Y \leq y_2) - P(X \leq x_1, y_1 < Y \leq y_2) \\
&= P(x_1 < X \leq x_2, y_1 < Y \leq y_2) \geq 0 \quad (P(\dots) \in [0, 1]!) \quad \text{☺}
\end{aligned}$$

Randverteilungen, unabhängige Zufallsvariablen — Répartitions marginales, variables aléatoires indépendantes

Für die Restriktionen definieren wir: • *Quant aux restrictions nous définissons:*

Definition: • **Définition:** **Randverteilungen** der zweidimensionalen Zufallsvariablen $\vec{X}$:
 • **distributions (répartitions) marginales de la variable aléatoire bidimensionnelle $\vec{X}$:**

$$F_X(x) := F(x, \infty) = P(X \leq x) = P(X \leq x, Y \leq \infty)$$

$$F_Y(y) := F(\infty, y) = P(Y \leq y) = P(X \leq \infty, Y \leq y)$$

Bemerkung: • **Remarque:** $F_X(x)$ und $F_Y(y)$ sind die Verteilungsfunktionen der Komponenten X und Y .
 • $F_X(x)$ et $F_Y(y)$ sont les fonctions des composantes X et Y .

Sei • *Soit* $A = \{X \leq x, Y \leq \infty\}$, $B = \{X \leq \infty, Y \leq y\} \Rightarrow A \cap B = \{X \leq x, Y \leq y\}$

Für unabhängige Ereignisse gilt nun: • *Pour des événements indépendants il vaut:*

$$P(A \cap B) = P(A) \cdot P(B)$$

$$\begin{aligned}
\leadsto F(x, y) &= P(X \leq x, Y \leq y) = P(A \cap B) = P(A) \cdot P(B) \\
&= P(X \leq x, Y \leq \infty) \cdot P(X \leq \infty, Y \leq y) = F_X(x) \cdot F_Y(y)
\end{aligned}$$

Daher können wir den Begriff der Unabhängigkeit von den Ereignissen auf die Variablen übertragen:

• *Par conséquent nous pouvons transférer la notion d'indépendance des événements aux variables:*

Definition: • **Définition:** X und Y heissen **unabhängig**, wenn gilt:
 • X et Y s'appellent **indépendantes**, s'il vaut:
 $\forall_{(x,y)} F(x, y) = F_X(x) \cdot F_Y(y)$

Identisch verteilte Variablen — Variables réparties de façon identique

Sei • *Soit* $\vec{X} = \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix}$, $\vec{X}^T = (X_1, \dots, X_n) \leadsto F_{\vec{X}}(x_1, x_2, \dots, x_n) = P(X_1 \leq x_1, \dots, X_n \leq x_n)$

Definition: • **Définition:** Falls alle Vektorkomponenten X_i , ($i = 1, \dots, n$) dieselbe Verteilungsfunktion $F_Z(Z)$, ($Z = X_1, \dots, X_n$) besitzen, so heissen die X_i **identisch verteilt**.
 • *Si toutes les composantes X_i , ($i = 1, \dots, n$) ont la même fonction de répartition $F_Z(Z)$, ($Z = X_1, \dots, X_n$), les X_i s'appellent **répartis de façon identique**.*

Definition: • **Définition:** Die Variablen X_i , ($i = 1, \dots, n$) heissen **vollständig unabhängig**, wenn jede der Variablen von allen andern unabhängig ist.
 • *Les variables X_i , ($i = 1, \dots, n$) s'appellent **complètement indépendantes**, si chacune des variables est indépendante de toutes les autres.*

Mittels der Definition für $n = 2$ kann man folgern: • *A l'aide de la définition pour $n = 2$ on peut déduire:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

X_i , ($i = 1, \dots, n$) vollständig unabhängig und identisch verteilt
 • X_i , ($i = 1, \dots, n$) *compl. indépendants et réparties de façon identique*

Beh.: • **Thè.:**

$$\forall_{(x_1, x_2, \dots, x_n)} F_{\vec{X}}(x_1, x_2, \dots, x_n) = F(x_1) \cdot F(x_2) \cdot \dots \cdot F(x_n)$$

Bemerkung: • **Remarque:** F bestimmt in diesem Fall also $F_{\vec{X}}$.
 • *Dans ce cas, c'est F qui détermine $F_{\vec{X}}$.*

4.10.2 Der diskrete Fall — Le cas discret

Sei • *Soit $p_{ik} := P(X = x_i, Y = y_k) := f(x_i, y_k)$, in $\mathbb{N}$*

Allgemein sagen wir: • *Généralement nous disons:*

Definition: • **Définition:** $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ heisst **diskret**, wenn $\{x_i, y_k\}$ resp. $\{x_i, y_k, p_{ik}\}$ abzählbar ist.
 • $\vec{X} = \begin{pmatrix} X \\ Y \end{pmatrix}$ s'appelle **discret**, si $\{x_i, y_k\}$ resp. $\{x_i, y_k, p_{ik}\}$ est dénombrable.

Formel: • **Formule:** $F(x, y) = \sum_{x_i \leq x, y_k \leq y} p_{ik}$

Bsp.: • **Exemple:**

Bei der Qualitätskontrolle eines Loses von Zahnradwellen werden die Dicke d und gleichzeitig am selben Stück jeweils auch die Länge l kontrolliert. (Dicke und Länge werden jeweils in einem separaten Arbeitsgang gefertigt. Sie können daher als unabhängig angesehen werden.) Wir können die Werte von X und Y frei wie folgt festlegen:

• *Au contrôle de qualité d'un lot d'arbres de roues dentées on contrôle l'épaisseur d et simultanément à la même pièce chaque fois aussi la longueur l . (L'épaisseur et la longueur sont fabriquées chaque fois pendant une phase de travail séparée. Par conséquent on les peut considérer comme indépendents.) Nous pouvons définir librement les valeurs de X et de Y comme il suit:*

Kriterium: • Critère:	X	Y
d innerhalb der Toleranz • d dans la tolérance	$x_1 = 0$	—
d ausserhalb Toleranz • d hors de la tolérance	$x_2 = 1$	—
l innerhalb der Toleranz • l dans la tolérance	—	$y_1 = 0$
l ausserhalb Toleranz • l hors de la tolérance	—	$y_2 = 1$

Aus Erfahrung wissen wir: • *Nous savons par l'expérience:*

Kriterium: • Critère:	Menge: • Montant:
Ausschuss • <i>Rebut</i>	5%
<i>d</i> falsch • <i>d fautif</i>	1%
<i>l</i> falsch • <i>l fautif</i>	3%
<i>d</i> und <i>l</i> falsch • <i>d et l fautifs</i>	1%

Konsequenz: • Conséquence:

$X \backslash Y$	0	1	
0	$p_{11} = 0.95$	$p_{21} = 0.01$	$p_{\cdot 1} = 0.96$
1	$p_{12} = 0.03$	$p_{22} = 0.01$	$p_{\cdot 2} = 0.04$
	$p_{1\cdot} = 0.98$	$p_{2\cdot} = 0.02$	$p_{tot} = 1$

(Z.B. ist: • *P.ex. on a:* $p_{11} + p_{21} = p_{\cdot 1}$, $p_{\cdot 1} = P(Y = 0) \dots$)

Definition: • **Définition:** Wahrscheinlichkeiten wie $f_Y(Y = 1) := p_{\cdot 1}$ heissen **Randsummen**.

• *Probabilités comme $p_{\cdot 1}$ s'appellent sommes marginales.*

Übersicht über die Randsummen: • *Vue d'ensemble des sommes marginales:*

$$\begin{aligned}
 p_{\cdot 1} &= P(Y = 0) = p_{11} + p_{21} = 0.96 \\
 p_{\cdot 2} &= P(Y = 1) = p_{12} + p_{22} = 0.04 \\
 p_{1\cdot} &= P(X = 0) = p_{11} + p_{12} = 0.98 \\
 p_{2\cdot} &= P(X = 1) = p_{21} + p_{22} = 0.02
 \end{aligned}$$

Es gilt: • *Il vaut:*

Satz: • **Théorème:**

$$\begin{aligned}
 f_X(X = x_i) &= p_{i\cdot} = P(X = x_i) = \sum_k p_{ik} \\
 f_Y(Y = y_i) &= p_{\cdot k} = P(Y = y_k) = \sum_i p_{ik} \\
 \sum_k p_{\cdot k} &= p_{\cdot i} p_{i\cdot} = 1
 \end{aligned}$$

Aus der allgemeinen Definition der Unabhängigkeit von Variablen folgt:

• *De la définition générale de l'indépendance de variables il suit:*

Korollar: • **Corollaire:** Diskrete Variablen X und Y unabhängig
 • *Variables discrètes X et Y indépendantes*

$$\forall_{(i,k)} p_{ik} = p_{i\cdot} \cdot p_{\cdot k}$$

Bsp.: • **Exemple:** $p_{1\cdot} \cdot p_{\cdot 1} = 0.98 \cdot 0.96 \approx 0.94 \neq p_{11} = 0.95$

$\leadsto$ In diesem Beispiel sind die Variablen **nicht unabhängig!** $\leadsto$ dans cet exemple les variables **ne sont pas indépendantes!**

4.10.3 Der stetige Fall — Le cas continu

Wir betrachten wieder den zweidimensionalen Modellfall.

- *Nous considérons de nouveau le cas modèle bidimensionnel.*

Definition: • **Définition:**

Wir nennen $\vec{X}$ stetig, wenn $\vec{X}$ überabzählbar-unendlich viele reelle Wertepaare annehmen kann und zudem eine **Wahrscheinlichkeitsdichte** $f(x, y)$ gegeben ist mit einer (stückweise) stetigen Verteilungsfunktion $F(x, y)$ mit:

- *Nous appelons $\vec{X}$ continu, si $\vec{X}$ peut atteindre des paires de valeurs réelles plus que dénombrables de façon infinie et si en plus une **densité aléatoire** $f(x, y)$ est donnée avec une fonction de répartition $F(x, y)$ qui est continue (par morceaux):*

$$F(x, y) = \int_{-\infty}^x \int_{-\infty}^y f(u, v) du dv$$

Aus den allgemeinen Definitionen folgt: • *On déduit des définitions générales:*

Folgerung: • **Conclusion:**

1. $f(x, y) \geq 0$
2. $F(\infty, \infty) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f(u, v) du dv = 1$

Für die **Dichte der Randverteilungen** gilt: • *Pour la densité des **répartitions marginales** il vaut:*

Folgerung: • **Conclusion:**

$$f_X(x) = \int_{-\infty}^{\infty} f(x, y) dy, \quad f_Y(y) = \int_{-\infty}^{\infty} f(x, y) dx$$

Aus der Unabhängigkeitsdefinition durch die Verteilungsfunktion ($F(x, y) = F_X(x) \cdot F_Y(y)$) kann man mit Hilfe des Mittelwertsatzes der Integralrechnung folgern:

- *De la définition de l'indépendance par la fonction de répartition ($F(x, y) = F_X(x) \cdot F_Y(y)$) on peut conclure à l'aide du théorème de la moyenne du calcul intégral:*

Satz: • **Théorème:**

Vor.: • **Hyp.:**

- X, Y unabhängig
- X, Y indépendantes

Beh.: • **Thè.:**

$$f(x, y) = f(x) \cdot f(y)$$

4.11 Mehrdimensionale Erwartung — Espérance multidimensionnelle

4.11.1 Erwartung, Mittelwert — Espérance, moyenne

Wir betrachten wieder den Modellfall $n = 2$. Sei $f(x, y)$ die Wahrscheinlichkeitsfunktion resp. die Dichte. Analog zum eindimensionalen Fall kann man auch jetzt Kenngrößen für X oder Y definieren: • *Nous considérons de nouveau le cas modèle $n = 2$. Soit $f(x, y)$ la fonction aléatoire resp. la densité. Analogiquement au cas unidimensionnel on peut maintenant aussi définir des grandeurs caractéristiques pour X ou Y :*

Definition: • **Définition:** Erwartungswert $E(g(X, Y))$ einer gegebenen Funktion $g(X, Y)$:
• **Valeur d'espérance** $E(g(X, Y))$ d'une fonction donnée $g(X, Y)$:

$$E(g(X, Y)) : = \begin{cases} \sum_i \sum_k g(x_i, y_k) \cdot f(i, k) & \text{dis.} \\ \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} g(x, y) \cdot f(x, y) dx dy & \text{cont.} \end{cases}$$

Voraussetzung: • *Prémisse ou hypothèse:*

$$\sum_i \sum_k |g| \cdot f \text{ oder } \bullet \text{ ou } \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} |g| \cdot f dx dy$$

exist. resp. conv. • *exist. resp. conv.*

Wie im eindimensionalen Fall folgert man (Rechnung):

• *Analogiquement au cas unidimensionnel on déduit (calcul):*

Satz: • **Théorème:** $E(a g(X, Y) + b h(X, Y)) = a E(g(X, Y)) + b E(h(X, Y))$

Speziell: • **Spécialement:** $a = b = 1$, $g(X, Y) = X$, $h(X, Y) = Y \rightsquigarrow$

Korollar: • **Corollaire:** $E(X + Y) = E(X) + E(Y)$

Verallgemeinerung: • *Généralisation:*

Satz: • **Théorème:** Additionssatz für Mittelwerte
• **Théorème d'addition pour les moyennes**

$$E(X_1 + X_2 + \dots + X_n) = E(X_1) + E(X_2) + \dots + E(X_n)$$

Allgemein ist $E(X^2) \neq E(X)^2$. • *Généralement il est $E(X^2) \neq E(X)^2$.*

Bsp.: • **Exemple:** Würfeln • *Jouer au dés* $\rightsquigarrow E(X) = \frac{7}{2}$
$$E(X^2) = \sum_{k=1}^6 k^2 \cdot \frac{1}{6} = \frac{91}{6} \neq \left(\frac{7}{2}\right)^2 = E(X)^2$$

Für unabhängige Variablen gilt jedoch infolge von $f(x, y) = f(x) \cdot f(y)$ resp. von $F(x, y) = F(x) \cdot F(y)$ z.B. im diskreten Fall:

• Pour des variables indépendantes il vaut, à cause de la formule $f(x, y) = f(x) \cdot f(y)$ resp. à cause de la formule $F(x, y) = F(x) \cdot F(y)$, p.ex. dans le cas discret:

$$E(XY) = \sum_i \sum_k x_i \cdot y_k \cdot f(x_i, y_k) = \sum_i x_i \cdot f_X(x_i) \cdot \sum_k y_k \cdot f_Y(y_k) = E(X) \cdot E(Y)$$

Entsprechend im stetigen Fall. • Correspondamment au cas continu. $\leadsto$

Konsequenz: • **Conséquence:**

X, Y unabhängig • X, Y indépendantes $\Rightarrow E(XY) = E(X) \cdot E(Y)$

Verallgemeinerung: • *Généralisation:*

Satz: • **Théorème:** **Multiplikationssatz für Mittelwerte**
 • **Théorème de multiplications pour les moyennes**

Vor.: • **Hyp.:**

$i \neq k \Rightarrow X_i, X_k$ unabhängig • indépendantes

Beh.: • **Thè.:**

$$E(X_1 \cdot X_2 \cdot \dots \cdot X_n) = E(X_1) \cdot E(X_2) \cdot \dots \cdot E(X_n)$$

4.11.2 Varianz, Kovarianz, Korrelation — Variance, covariance, corrélation

Bekannt vom Fall mit einer Variablen: • *Connu du cas avec une variable:*

$$\sigma^2 = E(X^2) - \mu^2 = E((X - \mu_X)^2)$$

Wir definieren: • *Nous définissons:*

Definition: • **Définition:** **Varianzen:** • **Variances:**
 $\sigma_X^2 = \text{Var}(X) = E((X - \mu_X)^2)$
 $\sigma_Y^2 = \text{Var}(Y) = E((Y - \mu_Y)^2)$

Sei • *Soit* $Z = X + Y$

$$\begin{aligned} \leadsto \sigma_{X+Y}^2 &= \sigma_Z^2 = E(Z^2) - (\mu_Z)^2 = E(Z^2) - (E(Z))^2 = E(X^2 + 2XY + Y^2) - (E(X + Y))^2 \\ &= E(X^2) + 2E(XY) + E(Y^2) - (E(X) + E(Y))^2 \\ &= E(X^2) + 2E(XY) + E(Y^2) - E(X)^2 - 2E(X)E(Y) - E(Y)^2 \\ &= (E(X^2) - E(X)^2) + 2(E(XY) - E(X)E(Y)) + (E(Y^2) - E(Y)^2) \\ &= \sigma_X^2 + 2(E(XY) - E(X)E(Y)) + \sigma_Y^2 \end{aligned}$$

Definition: • **Définition:** $\sigma_{XY} := \text{cov}(XY) := (E(XY) - E(X)E(Y))$
 heisst **Kovarianz** von X und Y
 • s'appelle **covariance** de X et Y

Es gilt: • *Il vaut:*

$$\begin{aligned}
E((X - \mu_X)(Y - \mu_Y)) &= E(X \cdot Y - X \cdot \mu_Y - \mu_X \cdot Y + \mu_X \cdot \mu_Y) \\
&= E(X \cdot Y) - \mu_Y \cdot E(X) - \mu_X \cdot E(Y) + \mu_X \cdot \mu_Y \cdot E(1) = E(X \cdot Y) - \mu_Y \cdot \mu_X - \mu_X \cdot \mu_Y + \mu_X \cdot \mu_Y \\
&= E(X \cdot Y) - \mu_X \cdot \mu_Y = E(X \cdot Y) - E(X)E(Y) \Rightarrow E((X - \mu_X)(Y - \mu_Y)) = \sigma_{XY}
\end{aligned}$$

Definition: • **Définition:** $\rho_{XY} := \frac{E((X - \mu_X)(Y - \mu_Y))}{\sigma_X \sigma_Y} = \frac{\sigma_{XY}}{\sigma_X \sigma_Y} = \frac{\text{cov}(XY)}{\sigma_X \sigma_Y}$
 heisst **Korrelationskoeffizient** von X und Y
 • s'appelle **coefficient de corrélation** de X et Y
 ($\sigma_X, \sigma_Y \neq 0$)

Man sieht unmittelbar: • Il est évident:

Satz: • **Théorème:**

1. $\sigma_{X+Y}^2 = \sigma_Z^2 = \sigma_X^2 + \sigma_Y^2 + 2\sigma_{XY} = \sigma_X^2 + \sigma_Y^2 + 2\text{cov}(XY)$
2. $\sigma_{XY} = \text{cov}(XY) = E(XY) - \mu_X \mu_Y$
 $= E(X \cdot Y) - E(X)E(Y) = E((X - \mu_X)(Y - \mu_Y))$
3. $\sigma_{XY} = \sigma_{YX}, \text{cov}(XY) = \text{cov}(YX)$
4. $X = Y \Rightarrow |\rho_{XY}| = 1$
5. $|\rho_{XY}| \leq 1$

Zur letzten Behauptung (die andern sind bewiesen) betrachten wir als Beispiel den stetigen Fall:

• Quant au dernier théorème (les autres ont été vérifiés) nous considérons comme exemple le cas continu:

$$\begin{aligned}
\rho_{XY} &:= \frac{E((X - \mu_X)(Y - \mu_Y))}{\sigma_X \sigma_Y} = \frac{E((X - \mu_X)(Y - \mu_Y))}{(E((X - \mu_X)^2) E((Y - \mu_Y)^2))^{\frac{1}{2}}} \\
&\Rightarrow \rho_{XY} \cdot (E((X - \mu_X)^2) E((Y - \mu_Y)^2))^{\frac{1}{2}} = E((X - \mu_X)(Y - \mu_Y)) \\
&\Rightarrow \rho_{XY} \cdot \underbrace{\left(\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \mu_X)^2 \cdot f(x, y) dx dy \right)^{\frac{1}{2}}}_{I_1} \cdot \underbrace{\left(\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (y - \mu_Y)^2 \cdot f(x, y) dx dy \right)^{\frac{1}{2}}}_{I_2} = \rho_{XY} \cdot (I_1 \cdot I_2)^{\frac{1}{2}} \\
&= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \mu_X)(y - \mu_Y) \cdot f(x, y) dx dy \Rightarrow \rho_{XY} \cdot (I_1 \cdot I_2)^{\frac{1}{2}} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \mu_X)(y - \mu_Y) \cdot f(x, y) dx dy
\end{aligned}$$

Hier wenden wir auf $I_1 \cdot I_2$ die Ungleichung von Cauchy-Schwarz an.

• Ici nous appliquons l'inégalité de Cauchy-Schwarz à $I_1 \cdot I_2$.

$$(\text{Cauchy-Schwarz: } (\int_G |f g|)^2 \leq (\int_G f^2) \cdot (\int_G g^2), \quad \int_G f g \leq \int_G |f g| \Rightarrow (\int_G f g)^2 \leq (\int_G f^2) \cdot (\int_G g^2))$$

(Die Ungleichung von Cauchy-Schwarz ist bekannt als „Skalarprodukt-Ungleichung“. Sie gilt auch im diskreten Fall entsprechend.)

• (L'inégalité de Cauchy-Schwarz est connue comme "inégalité du produit scalaire". Elle est correspondamment aussi valable dans le cas discret.)

Daraus folgert man: • On en déduit:

$$\rho_{XY} \cdot \left(\int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \mu_X) \cdot (y - \mu_Y) \cdot f(x, y) dx dy \right)^{\frac{1}{2}} \leq \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x - \mu_X)(y - \mu_Y) \cdot f(x, y) dx dy \Rightarrow |\rho_{XY}| \leq 1$$



Seien X und Y unabhängig • X et Y soient indépendantes $\leadsto P(XY) = P(X) \cdot P(Y) \Rightarrow \text{cov}(XY) = \sigma_{XY} = E(XY) - E(X)E(Y) = E(X)E(Y) - E(X)E(Y) = 0 \Rightarrow \rho_{XY} = \frac{\text{cov}(XY)}{\sigma_X \sigma_Y} = 0, (\sigma_X, \sigma_Y \neq 0)$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

X und Y unabhängig
• X et Y indépendantes

Beh.: • **Thè.:**

$\text{cov}(XY) = \rho_{XY} = 0$

Bemerkung: • **Remarque:**

Im Falle $\text{cov}(XY) = E((X - \mu_X)(Y - \mu_Y)) = 0$ kann man $E((X - \mu_X)(Y - \mu_Y))$ als ein Skalarprodukt begreifen, dass 0 ist. Man hat es daher mit senkrecht stehenden Vektoren zu tun.

Im Falle $\rho_{XY} = 1$ wird die oben angewandte Cauchy-Schwarzsche Ungleichung zu einer Gleichung. In diesem Fall kann man den folgenden Satz folgern (Beweis vgl. Kreyszig, Bibl. A10):

• Dans le cas $\text{cov}(XY) = E((X - \mu_X)(Y - \mu_Y)) = 0$ on peut interpréter $E((X - \mu_X)(Y - \mu_Y))$ comme produit scalaire qui est 0. On l'a donc à faire avec des vecteurs perpendiculaires.

Dans le cas $\rho_{XY} = 1$ l'inéquation de Cauchy-Schwarz qu'on vient d'appliquer en haut, devient une équation. Dans ce cas on peut déduire: (Preuve voir Kreyszig, Bibl. A10):

Satz: • **Théorème:**

Zwischen zwei Zufallsvariablen X und Y mit positiver Varianz besteht dann und nur dann eine lineare Beziehung $Y = \alpha X + \beta$ resp. $X = \gamma Y + \delta$, wenn $|\rho_{XY}| = 1$ gilt.

• Entre les variables aléatoires X et Y à des variances positives on a une relation linéaire $Y = \alpha X + \beta$ resp. $X = \gamma Y + \delta$ exactement si'il est $|\rho_{XY}| = 1$.

Konsequenz: • **Conséquence:**

Der Korrelationskoeffizient (vgl. auch Seite 227) ist somit ein **Mass für den linearen Zusammenhang zweier Zufallsvariablen**. Dieser Sachverhalt kann zur Prüfung eines linearen Zusammenhanges benutzt werden.

• Le coefficient de corrélation (voir aussi page 227) est par conséquent une **mesure pour le rapport linéaire de deux variables de hasard**. Ce fait peut être utilisé pour l'examen d'un rapport linéaire.

Achtung: • **Attention:**

Ein formaler Zusammenhang beweist keinen kausalen Zusammenhang!

• Un rapport formel ne prouve pas un rapport causal.

Bsp.: • **Exemple:** (Beispiel zu „formalem und kausalem Zusammenhang“.) Von Norden bis Süden nimmt die Körpergröße h ab, jedoch die Anzahl der Katholiken k zu. Die Korrelation $h \sim \frac{1}{k}$ ist aber bekanntlich eine **Scheinkorrelation!**

• (Un exemple quant à "rapport formel et rapport causal".) Plus on va vers le sud et plus la taille des

gens h diminue, cependant le nombre des catholiques k augmente. La corrélation $h \sim \frac{1}{k}$ est, comme chacun sait, une **corrélation fictive (feinte, fallacieuse, fausse)!**

Bemerkung: • **Remarque:** Scheinkorrelation findet man häufig in interessenpolitischen Argumenten. Für den Laien sind sie oft plausibel. . .
 • *Souvent, dans des arguments politiques qui se vent à des intérêts privés, on trouve des corrélation fausses. Pour le laïc (l'ignorant, l'amateur) ces arguments semblent être logiques. . .*

Sei • *Soit* $X + Y = Z$, X, Y unkorreliert • *non corrélées* $\leadsto$
 $\text{cov}(X, Y) = 0 = E(XY) - E(X)E(Y) \Rightarrow E(XY) = E(X)E(Y) \leadsto X, Y$ unabhängig • *indépendantes*

Oben haben wir hergeleitet: • *En haut on a déduit:*

$$\sigma_Z^2 = \sigma_X^2 + \sigma_Y^2 + 2\sigma_{XY} = \sigma_X^2 + \sigma_Y^2 + 2\text{cov}(XY) + \sigma_X^2 + \sigma_Y^2 + 0 \Rightarrow \text{Var}(X + Y) = \text{Var}(X) + \text{Var}(Y)$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$X_1, \dots, X_n$ unkorreliert • *non corrélées*

Beh.: • **Thè.:**

1. $X_1, \dots, X_n$ unabhängig • *indépendantes*
2. $\text{Var}(X_1 + \dots + X_n) = \text{Var}(X_1) + \dots + \text{Var}(X_n)$

4.11.3 Der diskrete Fall — Le cas discret

Folgerung: • **Conclusion:**

1. $E(g(X, Y)) = E(g(X))$
 $\Rightarrow E(g(X)) = \sum_i g(x_i) \sum_k f(x_i, y_k) = \sum_i g(x_i) \cdot f_X(x_k)$
 $\leadsto$ Formel für eine einzige Variabel, o.k.!
 • *Formule pour une seule variable, o.k.!*
2. **Erwartungswerte:** • **Valeurs d'espérance.**
 $\mu_X := E(X) := \sum_i x_i \cdot f_X(x_i) = \sum_i \sum_k x_i \cdot p_{ik}$
 $\mu_Y := E(Y) := \sum_k y_k \cdot f_X(y_k) = \sum_k \sum_i y_k \cdot p_{ik}$

4.11.4 Der stetige Fall — Le cas continu

Analog zum diskreten Fall kann man auch jetzt Kenngrößen für X oder Y definieren:

• *Analogiquement au cas discret on peut maintenant définir des grandeurs caractéristiques pour X ou Y :*

Definition: • **Définition:**

1. $E(g(X, Y)) = E(g(X))$
 $\Rightarrow E(g(X)) = \int_{-\infty}^{\infty} g(x) \int_{-\infty}^{\infty} f(x, y) dy dx = \int_{-\infty}^{\infty} g(x) \cdot f_X(x) dx$
 $\leadsto$ Formel für eine einzige Variabel, o.k.!
 • *Formule pour une seule variable, o.k.!*

2. **Erwartungswerte:** • **Valeurs d'espérance.**

$$\mu_X := E(X) := \int_{-\infty}^{\infty} x \cdot f_X(x) dx = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x \cdot f(x, y) dx dy$$

$$\mu_Y := E(Y) := \int_{-\infty}^{\infty} y \cdot f_Y(y) dy = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} y \cdot f(x, y) dx dy$$

4.12 Mehrdimensionale Verteilungen — Répartitions multidimensionnelles

4.12.1 Zweidimensionale Normalverteilung — Distribution normale bidimensionnelle

Im stetigen Fall spielt diese Verteilung eine besondere Rolle.

- *Dans le cas continu cette distribution joue un rôle important.*

Definition: • **Définition:**

Gegeben seien die Parameter $\mu_X, \mu_Y, \sigma_X, \sigma_Y, \rho_{XY}$. $\begin{pmatrix} X \\ Y \end{pmatrix}$ besitzt eine zweidimensionale Normalverteilung, wenn die Dichte durch folgende Funktion gegeben ist:

- *Soient donnés les paramètres $\mu_X, \mu_Y, \sigma_X, \sigma_Y, \rho_{XY}$. $\begin{pmatrix} X \\ Y \end{pmatrix}$ possède une répartition normale bidimensionnelle si la densité est donnée par la fonction suivante:*

$$f(x, y) = \frac{1}{2\pi\sigma_X\sigma_Y\sqrt{1-\rho_{XY}^2}} \cdot e^{-\frac{h(x, y)}{2}}$$

$$h(x, y) = \frac{1}{(1-\rho_{XY}^2)} \left(\frac{(x-\mu_X)^2}{\sigma_X^2} - 2\frac{\rho_{XY}(x-\mu_X)(y-\mu_Y)}{\sigma_X\sigma_Y} + \frac{(y-\mu_Y)^2}{\sigma_Y^2} \right)$$

$x, y, \in \mathbb{R}$

Ohne Beweis folgender Satz: • *Le théorème suivant sans preuve:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

1. $\vec{X}$ besitzt eine zweidimensionale Normalverteilung
• $\vec{X}$ a une distribution normale bidimensionnelle
2. $\text{conv}(X Y) = 0$

Beh.: • **Thè.:**

1. X, Y unabhängig • X, Y indépendantes
2. $f(x, y) = \frac{1}{2\pi\sigma_X\sigma_Y} \cdot e^{-\frac{1}{2}\left(\frac{(x-\mu_X)^2}{\sigma_X^2} + \frac{(y-\mu_Y)^2}{\sigma_Y^2}\right)}$
3. $\mu_X = \mu_Y = 0, \sigma_X^2 = \sigma_Y^2 = 1 \Rightarrow f(x, y) = \frac{1}{2\pi} \cdot e^{-\frac{1}{2}(x^2+y^2)}$

4.12.2 Stichprobenfunktion, Testverteilungen — Fonctions d'échantillonnage, distributions de test

Bemerkung: • **Remarque:**

Die in der Statistik wichtigen Verteilungen kann man in zwei Klassen einteilen: Verteilungen, die im Zusammenhang von **mathematischen Modellen** von Zufallsexperimenten auftreten und Verteilungen, **Testverteilungen** oder **Prüfverteilungen**, die die Grundlage statistischer Tests bilden. Testverteilungen sind hier Verteilungen von Stichprobenfunktionen.

• *On peut diviser les distributions importantes des statistiques en deux classes: Les répartitions qui apparaissent en rapport avec les **modèles mathématiques** d'expériences aléatoires et répartitions, qui sont des **distributions de test** ou des **lois usuelles pour les tests**, ce qui forme la base de tests statistiques. Les distributions de test sont ici des distributions de fonctions d'échantillonnage.*

Vor.: • **Hyp.:**

Seien im Folgenden alle Komponenten $X_i, i = 1, \dots, n$ von $\vec{X}$ unabhängig und identisch normalverteilt mit dem Erwartungswert μ und der Varianz σ^2 . ($\leadsto$ Stichprobe vom Umfang n mit $N(\mu, \sigma^2)$. Die Grundgesamtheit sei normalverteilt.)

• *Dans ce qui suit toutes les composantes $X_i, i = 1, \dots, n$ de $\vec{X}$ soient indépendamment et identiquement distribuées de façon normale avec la valeur d'espérance μ et la variance σ^2 . ($\leadsto$ Echantillon de la taille n avec $N(\mu, \sigma^2)$. L'ensemble de base soit distribué de façon normale.)*

Mittelwert-Verteilung — Distribution de la moyenne

Diese Verteilung ist die Verteilung der Stichprobenfunktion $\bar{X}$. Sie ist wichtig bei der Behandlung der Vertrauensintervalle (**Konfidenzintervalle**) für den Mittelwert einer Normalverteilung bei bekannter Varianz. Dabei gilt:

• *Cette répartition est une distribution de la fonction d'échantillonnage $\bar{X}$. Elle est importante lors du*

traitement des **intervalles de confiance** pour la moyenne d'une distribution normale à la variance connue. Il vaut:

$$\bar{X} = \frac{1}{n} (X_1 + X_2 + \dots, X_n) = \frac{1}{n} \sum_{i=1}^n X_i$$

Satz: • **Théorème:** $\bar{X}$ ist normalverteilt mit dem Erwartungswert μ und der Varianz $\frac{\sigma^2}{n}$
 • $\bar{X}$ est distribué de façon normale avec la valeur d'espérance μ et la variance $\frac{\sigma^2}{n} \rightsquigarrow N(\mu, \frac{\sigma^2}{n})$

Bemerkung: • **Remarque:** Es ist wichtig zu bemerken, dass der Satz auch näherungsweise noch gilt, wenn die dazugehörige Grundgesamtheit mit denselben Parametern beliebig verteilt ist.
 • Il est important de remarquer que le théorème vaut encore de façon approximative si la base des données est distribuée de façon quelconque avec les mêmes paramètres.

Sei • Soit $\bar{X} \xrightarrow{h} Z = h(\bar{X}) = \frac{\bar{X} - \mu}{\sigma} \sqrt{n}$

Die Koordinatentransformation h bewirkt, dass der Satz ein Korollar ist aus folgendem Lemma:

• La transformation de coordonnées h cause que ce théorème est un corollaire du Lemme suivant:

Lemma: • **Lemme:** **Vor.:** • **Hyp.:**

$X_i, i = 1, \dots, n$ von • de $\vec{X}$
 X_i unabhängig, identisch normalverteilt mit $N(\mu, \sigma^2)$
 • X_i indépendantes, distribuées de façon normale avec $N(\mu, \sigma^2)$
 $\bar{X} = \frac{1}{n} (X_1 + X_2 + \dots, X_n)$
 $Z = \frac{\bar{X} - \mu}{\sigma} \sqrt{n}$

Beh.: • **Thè.:**

Z genügt der standardisierten Normalverteilung nach $N(0, 1)$
 • Z satisfait la ditribution normale standardisée d'après $N(0, 1)$

Dieses Lemma wiederum ist ein Korollar (Spezialfall) eines allgemeineren Satzes, wie man unmittelbar einsehen kann:

• Ce lemme est de nouveau un corollaire d'un théorème plus général comme on voit tout de suite:

Satz: • **Théorème:** **Summe unabhängiger normalverteilter Variablen**
 • **Somme de variables indépendantes et distribuées de façon normale**

Vor.: • **Hyp.:**

Seien $X_1, X_2, \dots, X_n$ unabhängig, normalverteilt mit

• *Soient $X_1, X_2, \dots, X_n$ indépendantes, distribuées de façon normale avec*

Mittelwerte • *Moyennes $\mu_1, \mu_2, \dots, \mu_n$*

Varianzen • *Variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$*

Beh.: • **Thè.:**

$\sum_{i=1}^n X_i$ normalverteilt • *distribuées de façon normale*

Mittelwert • *Moyenne $\mu = \sum_{i=1}^n \mu_i$*

Varianz • *Variance $\sigma^2 = \sum_{i=1}^n \sigma_i^2$*

Für den Beweis benötigen wir folgende Formel: • *Pour la preuve nous avons besoin de la formule suivante:*

Lemma: • **Lemme:**

Vor.: • **Hyp.:**

$Z = g(X, Y) = X + Y$ X, Y unabh. • *indep.*

Dichten: • *Densités:* $X \rightsquigarrow f_1(x), Y \rightsquigarrow f_2(y)$

Beh.: • **Thè.:**

$$F(Z) = \int_{-\infty}^{\infty} f_2(y) \left(\int_{-\infty}^{z-y} f_1(x) dx \right) dy$$

$$f_1 \in \text{stet'cont} \Rightarrow f(z) = \int_{-\infty}^{\infty} f_1(x) f_2(z-y) dx = \int_{-\infty}^{\infty} f_1(z-x) f_2(y) dy$$

Beweis Lemma: • **Preuve Lemme:**

Von Seite 132: • *De la page 132:* $f(x, y) = f_1(x) \cdot f_2(y)$

$$F(z) = \int_{x+y \leq z} f(x, y) dx dy = \int_{x+y \leq z} f_1(x) \cdot f_2(y) dx dy = \int_{-\infty}^{\infty} (f_2(y) \int_{-\infty}^{x=z-y} f_1(x) \cdot dx) dy \quad \text{☺}$$

$$z = x + y, \quad \frac{dz}{dx} = 1 \Rightarrow F(z) = \int_{-\infty}^{\infty} (f_2(y) \int_{-\infty}^z f_1(z-y) \cdot dx) dz \Rightarrow f(z) := \frac{d}{dz} F(z) =$$

$$\frac{d}{dz} \int_{-\infty}^{\infty} (f_2(y) \int_{-\infty}^z f_1(z-y) dz) dy = \int_{-\infty}^{\infty} (f_2(y) \frac{d}{dz} \int_{-\infty}^z f_1(z-y) dz) dy = \int_{-\infty}^{\infty} (f_2(y) f_1(z-y)) dy$$

Zieht man bei der Berechnung von F statt f_2 die Funktion f_1 nach vorne, so folgt:

• *Si en calculant F on met f_2 au lieu de f_1 en devant, on obtient:* $f(z) = \int_{-\infty}^{\infty} f_1(x) f_2(z-y) dy \quad \text{☺}$

Beweis Satz: • **Preuve Théorème:**

1. Induktion, Verankerung: $n = 1$: $X = X_1$ ist normalverteilt.
 • *Induction, "ancrer": $n = 1$: $X = X_1$ distribuée de façon normale.*
2. Induktion, Vererbungsgesetz von $n = 1$ auf $n = 2$. • *Induktion, transmission de $n = 1$ à $n = 2$.*
 Sei zuerst: • *Soit d'abord: $X = X_1 + X_2$*

Nach Vor.: • *D'après hyp.:* $f_1(x) = \frac{1}{\sqrt{2\pi}\sigma_1} \cdot e^{-\frac{1}{2}\left(\frac{x-\mu_1}{\sigma_1}\right)^2}$, $f_2(x) = \frac{1}{\sqrt{2\pi}\sigma_2} \cdot e^{-\frac{1}{2}\left(\frac{x-\mu_2}{\sigma_2}\right)^2}$

$$\Rightarrow f(x) = \int_{-\infty}^{\infty} f_1(x-y) f_2(y) dy = \int_{-\infty}^{\infty} \frac{1}{\sqrt{2\pi}\sigma_1} \cdot e^{-\frac{1}{2}\left(\frac{x-y-\mu_1}{\sigma_1}\right)^2} \frac{1}{\sqrt{2\pi}\sigma_2} \cdot e^{-\frac{1}{2}\left(\frac{y-\mu_2}{\sigma_2}\right)^2} dy$$

$$= \frac{1}{2\pi\sigma_1\sigma_2} \cdot \int_{-\infty}^{\infty} e^{-\frac{1}{2}\left(\frac{x-y-\mu_1}{\sigma_1} + \frac{y-\mu_2}{\sigma_2}\right)^2} dy$$

Sei • *Soit* $\mu := \mu_1 + \mu_2$, $\sigma^2 := \sigma_1^2 + \sigma_2^2$, $V := -\frac{1}{2}\left(\frac{x-y-\mu_1}{\sigma_1} + \frac{y-\mu_2}{\sigma_2}\right)^2$

Sei • *Soit* $V_1 := \frac{\sigma}{\sigma_1\sigma_2}\left(y - \frac{\sigma_1^2\mu_2 + \sigma_2^2(x-\mu_1)}{\sigma^2}\right)$, $V_2 := \frac{x-\mu}{\sigma}$

Durch ausmultiplizieren, gleichnamig machen und umformen zeigt man, dass die folgende Identität gilt (viele Terme $\leadsto$ z.B. mit Mathematica!):

• *Par multiplication des termes, mettre au même dénominateur et remodeler on démontre que l'identité suivante est valable (un tas de termes $\leadsto$ p.ex. avec Mathematica!):*

$$V \equiv V_1^2 + V_2^2$$

V_2 ist unabhängig von y , $V_1 := \tau$ wird als neue Integrationsvariable verwendet:

• V_2 est indépendante de y , $V_1 := \tau$ est utilisée comme nouvelle variable d'intégration:

$$\leadsto f(x) = \frac{1}{2\pi\sigma_1\sigma_2} \cdot \int_{-\infty}^{\infty} e^{-\frac{1}{2}(V_1^2+V_2^2)} dy = \frac{1}{2\pi\sigma_1\sigma_2} \cdot e^{-\frac{1}{2}V_2^2} \cdot \int_{-\infty}^{\infty} e^{-\frac{1}{2}V_1^2} dy$$

$$= \frac{1}{2\pi\sigma_1\sigma_2} \cdot e^{-\frac{1}{2}V_2^2} \cdot \int_{-\infty}^{\infty} e^{-\frac{1}{2}\tau^2} d\tau$$

Es gilt: • *Il vaut:* $\frac{dV_1}{dy} = \frac{d\tau}{dy} = \frac{\sigma}{\sigma_1\sigma_2} \Rightarrow dy = \frac{\sigma_1\sigma_2}{\sigma} d\tau$, $\int_{-\infty}^{\infty} e^{-\frac{1}{2}\tau^2} d\tau = \sqrt{2\pi}$ (p. 116)

$$\Rightarrow f(x) = \frac{1}{2\pi\sigma_1\sigma_2} \cdot e^{-\frac{1}{2}V_2^2} \cdot \int_{-\infty}^{\infty} e^{-\frac{1}{2}\tau^2} \frac{\sigma_1\sigma_2}{\sigma} d\tau = \frac{1}{2\pi\sigma} \cdot e^{-\frac{1}{2}V_2^2} \cdot \sqrt{2\pi} = \frac{\sqrt{2\pi}}{\sigma} \cdot e^{-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2}$$

$\leadsto$ Für $n = 2$ ist X damit normalverteilt. • *Pour $n = 2$, X est donc distribuée de façon normale.*

3. Induktionsschluss von $n = m$ auf $n = m + 1$: • *Conclusion de l'induction de $n = m$ à $n = m + 1$:*

Vor.: • *Hyp.:* $Y_1 = X_{m+1}$, $Y_2 = \sum_{i=1}^m X_i$ normalverteilt. • *distribuées de façon normale.*

Gesetz: • *Loi:* Y_1, Y_2 normalverteilt • *distribuées de façon normale*

$\leadsto$ Wie bei X_1, X_2 auch Summe normalverteilt. • *Comme à X_1, X_2 : la somme est aussi distribuée de cette façon.*

$\leadsto$ Beh.: • *Thèse.:* $Y_1 + Y_2 = \sum_{i=1}^m X_1 + X_{m+1} = \sum_{i=1}^{m+1} X_i$ normalverteilt. • *distribuées de façon normale.*



Durch Integralsubstitution findet man noch den folgenden Satz: • *Par substitution sous l'intégrale on trouve le théorème suivant:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

X normalverteilt mit $N(\mu, \sigma^2)$
 • X distribuée de façon normale avec $N(\mu, \sigma^2)$
 $\tilde{X} = c_1 \cdot X + c_2, \quad c_1, c_2 = \text{const.}$ (Lineare Transf.) • (Transf. lin.)

Beh.: • **Thè.:**

$\tilde{X}$ normalverteilt • $\tilde{X}$ distribuée de façon normale
 $\tilde{\mu} = c_1 \cdot \mu + c_2, \quad \tilde{\sigma}^2 = c_1^2 \cdot \sigma^2$

Die Formeln für $\tilde{\mu}$ und $\tilde{\sigma}^2$ sind schon auf Seite 99 in allgemeiner Form hergeleitet worden.
 • *Les formules pour $\tilde{\mu}$ et $\tilde{\sigma}^2$ ont été déduites à la page 99 en forme plus générale.*

4.12.3 Chi-Quadrat-Verteilung — Distribution du Khi-deux

Definition, Verteilungsfunktion — **Definition, fonction de répartition**

Sei • *Soit* $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ (empirische Streuung) • *(variance empirique)*

Zur Erinnerung: • *Se rappeler:* $X_i \rightsquigarrow N(\mu, \sigma^2)$

Damit kann man die folgende Stichprobenfunktion bilden: • *A partir de ceci, on peut former la fonction d'échantillon suivante:*

Definition: • **Définition:** $\chi^2 = \frac{n-1}{\sigma^2} S^2 = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2$

Im Falle von $N(0, 1)$ ($\mu = 0, \sigma^2 = 1$) erhalten wir: Au cas de $N(0, 1)$ ($\mu = 0, \sigma^2 = 1$) nous obtenons:

$$\chi^2 = \sum_{i=1}^n X_i^2 = X_1^2 + X_2^2 + \dots + X_n^2$$

χ^2 ist wieder eine Zufallsvariable. Die Herleitung der **Wahrscheinlichkeitsdichte** von χ^2 ist ein etwas grösseres Unterfangen (vgl. z.B. Lit. Kreyszig, Bibl. A10). Aus Gründen des Umfangs soll hier auf die Wiedergabe verzichtet werden (vgl. Anhang). Man erhält:

• χ^2 est encore une variable de hasard. La déduction de la **densité de probabilité** de χ^2 demande beaucoup de temps et de place (voir, par exemple lit. Kreyszig, Bibl. A10). Pour des raisons de volume nous devons renoncer de reproduire la déduction ici (voir annexe). On obtient:

Satz: • **Théorème:** **Wahrscheinlichkeitsdichte** von χ^2
 • **Densité de probabilité** de χ^2

$$f(x) = \begin{cases} 0 & x \leq 0 \\ K_n x^{\frac{n-2}{2}} e^{-\frac{x}{2}} & x > 0 \end{cases}$$

$$K_n := \frac{1}{2^{\frac{n}{2}} \Gamma(\frac{n}{2})}$$

Die Verteilungsfunktion erhält man durch Integration: • *On obtient la fonction de distribution par intégration:*

Formel: • **Formule:**
$$F(x) = K_n \int_0^x u^{\frac{n-2}{2}} e^{-\frac{u}{2}} du$$

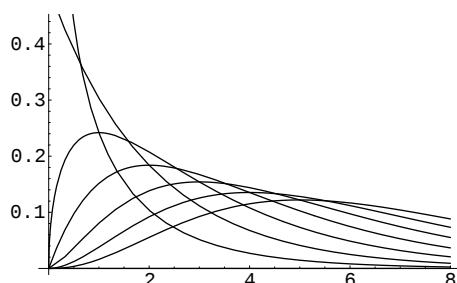
(K_n ist auf Grund der Forderung $\int_{-\infty}^{\infty} f(x) dx = \int_0^{\infty} f(x) dx = 1$ berechnet worden.)

• (K_n a été calculé sur la base de l'exigence $\int_{-\infty}^{\infty} f(x) dx = \int_0^{\infty} f(x) dx = 1$.)

Definition: • **Définition:**

n heisst die **Anzahl der Freiheitsgrade** der Verteilung. Γ ist die **Gammafunktion**.

• n s'appelle **nombre de degrés de liberté** de la distribution. Γ est la **fonction gamma**.



Das Bild zeigt $f(x)$ die für die Freiheitsgrade $n = 0, 1, \dots, 7$

• *L'image montre $f(x)$ pour les degrés de liberté $n = 0, 1, \dots, 7$*

Gammafunktion, Gamma- und Betaverteilung — Foncton gamma, répartition gamma et bêta

Für die Gammafunktion gilt: • *Pour la foncton gamma il vaut:*

1. $\Gamma(\alpha) := \int_0^{\infty} e^{-t} t^{\alpha-1} dt$

2. $\Gamma(\alpha + 1) = \alpha \Gamma(\alpha)$ (Partielle Integration) • *(Intégration partielle)*

3. $\Gamma(1) = \int_0^{\infty} e^{-t} t^{1-1} dt = \int_0^{\infty} e^{-t} dt = 1$

4. $\Gamma(2) = 1 \cdot \Gamma(1) = 1!$, $\Gamma(3) = 2 \cdot \Gamma(2) = 2!$, ..., $\Gamma(n+1) = n!$

5. $\Gamma(\frac{1}{2}) = \sqrt{\pi}$, $\Gamma(\frac{3}{2}) = \frac{1}{2} \Gamma(\frac{1}{2}) = \frac{\sqrt{\pi}}{2}$ u.s.w • *etc.*

Bemerkung: • **Remarque:**

Die Chi-Quadrat-Verteilung ist ein Sonderfall der **Gammaverteilung**.

• *La distribution du Khi-deux est un cas spécial de la distribution gamma.*

Definition: • **Définition:**Sei • *Soit* $\alpha > 0$

Die **Gammaverteilung** ist durch die folgende Dichtefunktion gegeben: • *La distribution gamma est donnée par la fonction de densité suivante:*

$$f(x) = \begin{cases} 0 & x < 0 \\ e^{-x} \frac{x^{\alpha-1}}{\Gamma(\alpha)} & x \geq 0 \end{cases}$$

Definition: • **Définition:**Sei • *Soit* $\alpha > 0, \beta > 0$

$B(\alpha, \beta)$ heisst **Beta-Funktion:** • $B(\alpha, \beta)$ s'appelle **fonction bêta:**

$$B(\alpha, \beta) := \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha + \beta)}$$

Definition: • **Définition:**Sei • *Soit* $\alpha > 0, \beta > 0$

Die **Betaverteilung** ist durch die folgende Dichtefunktion gegeben: • *La distribution bêta est donnée par la fonction de densité suivante:*

$$f(x) = \begin{cases} 0 & x \leq 0 \\ \frac{x^{\alpha-1} (1-x)^{\beta-1}}{B(\alpha, \beta)} & x > 0 \end{cases}$$

4.12.4 Sätze zur Chi-Quadrat-Verteilung — Théorèmes sur la distribution du Khi-deux

Mit Hilfe der Formeln für die Momente berechnet man:

• *A l'aide de la formule pour les moments on calcule:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Chi-Quadrat-Verteilung • *Distribution du Khi-deux*

Beh.: • **Thè.:**

$$1. \mu = n$$

$$2. \sigma^2 = 2n$$

Für grosse n kann man die Chi-Quadrat-Verteilung durch die Normalverteilung annähern. Es gilt:

• *Pour des n grands on peut approximer la distribution du Khi-deux par la distribution normale. Il vaut:*

Satz: • **Théorème:**

1. Die Zufallsvariable χ^2 ist asymptotisch normalverteilt mit $\mu = n$ und $\sigma^2 = 2n$. Für grosse n ist:
 - *La variable aléatoire χ^2 est répartie de façon normale asymptotique avec $\mu = n$ et $\sigma^2 = 2n$. Pour des grands n il vaut:*

$$F(x) \approx \Phi\left(\frac{x-n}{\sqrt{2n}}\right)$$

2. Die Zufallsvariable $\sqrt{2}\chi^2$ ist asymptotisch normalverteilt mit $\mu = \sqrt{2n-1}$ und $\sigma^2 = 1$. Für grosse n ist:
 - *La variable aléatoire $\sqrt{2}\chi^2$ est répartie de façon normale asymptotique avec $\mu = \sqrt{2n-1}$ et $\sigma^2 = 1$. Pour des grands n il vaut:*

$$F(x) \approx \Phi(\sqrt{2}x\sqrt{2n-1})$$

4.12.5 t-Verteilung von Student — Distribution de Student

Herleitung — Dédution

Die **t-Verteilung** bildet die Grundlage für wichtige Tests. Sie wurde von W.S. Gosset unter dem Decknamen „Student“ publiziert.

• *La distribution de Student est la base de tests importants. Elle a été publiée de W.S. Gosset sous le pseudonyme "Student".*

Wir bilden zu den voneinander unabhängigen Stichprobenfunktionen X (normalverteilt) und Y (χ_n^2 -verteilt) die neue Stichprobenfunktion t (oder T):

• *Quant aux fonctions d'échantillonnage indépendantes X (distribué de façon normale) et Y (distribué de façon χ_n^2) nous formons une nouvelle fonction d'échantillonnage t (ou T):*

$$T = \frac{X}{\sqrt{Y}} \sqrt{n}, \quad t = \frac{X}{\sqrt{Y}/\sqrt{n}}$$

Zur Erinnerung: • *Se rappeler:* $X_i \in N(\dots), i = 1, 2, \dots, n \rightsquigarrow \bar{X} \in N(\dots)$

Z.B. für $\mu = 0$ und $\sigma^2 = 1$ ($N(0, 1)$) erhält man: • *P.ex. pour $\mu = 0$ et $\sigma^2 = 1$ ($N(0, 1)$) on obtient:*

$$T = \frac{\bar{X}}{S} \sqrt{m} = \frac{\bar{X}}{\sqrt{Y/m}}, \quad Y = S^2, \quad m = n - 1 \quad (\text{Vgl. Vertrauensintervalle.} \quad \bullet \text{ Voir interv. de confiance.})$$

Definition: • **Définition:**

Die zu T resp. t gehörige Verteilung heisst **Student-Verteilung**.

• *La distribution liée à T resp. t s'appelle **distribution de Student**.*

Später werden wir zeigen: $Z = \bar{X}$ normalverteilt ($N(0, 1)$) und $Y = S^2$ ist $\chi_{m=n-1}^2$ -verteilt.

• *On va prouver plus tard: $Z = \bar{X}$ est distribuée de façon normale ($N(0, 1)$) et $Y = S^2$ est distribuée de façon $\chi_{m=n-1}^2$.*

Definition: • **Définition:**

n resp. m heisst die **Anzahl der Freiheitsgrade** der Verteilung.

• *n resp. m s'appelle **nombre de degrés de liberté** de la distribution.*

Die Herleitung der **Wahrscheinlichkeitsdichte** von T wollen wir auf später verschieben (vgl. z.B. Lit. Kreyszig, Bibl. A10 resp. Anhang). Man erhält ($T \rightsquigarrow z$):

• Nous faisons la déduction de la **densité de probabilité** de T plus tard (voir. par exemple lit. Kreyszig, Bibl. A10 ou annexe). On obtient ($T \rightsquigarrow z$):

Satz: • **Théorème:**

$$f(z) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \Gamma(\frac{n}{2})} \cdot \frac{1}{(1 + \frac{z^2}{n})^{(n+1)/2}}$$

$$F(z) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \Gamma(\frac{n}{2})} \cdot \int_{-\infty}^z \frac{1}{(1 + \frac{u^2}{n})^{(n+1)/2}} du$$

Definition: • **Définition:**

Für $n = 1$ erhalten wir die **Cauchy-Verteilung**.

• Pour $n = 1$ on obtient la **distribution de Cauchy**.

Sätze — Théorèmes

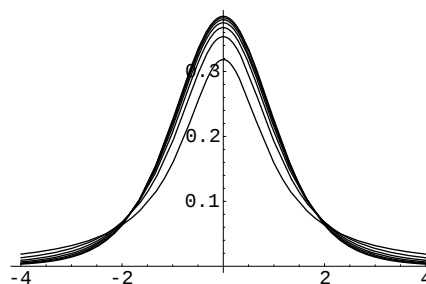
Formel: • **Formule:**

1. $n = 1$: Die Cauchy-Verteilung besitzt keinen Mittelwert und keine Varianz.
• $\rightsquigarrow$ La distribution de Cauchy n'a ni de moyenne ni de variance.
2. $n = 2 \rightsquigarrow$ keine Varianz • pas de variance
3. $n \geq 2$, $n \in \mathbb{N} \rightsquigarrow \mu = 0$ (z^2 , Symmetrie! • symétrie!)
4. $n \geq 3 \Rightarrow \sigma^2 = \frac{n}{n-2}$

Das Bild zeigt $f(x)$ die für die Freiheitsgrade $n = 1, \dots, 7$

• L'image montre $f(x)$ pour les degrés de liberté $n = 1, \dots, 7$

Es gilt: • Il vaut:



Satz: • **Théorème:**

Für $n \rightarrow \infty$ strebt die Verteilungsfunktion $f(x)$ ($z \rightsquigarrow x$) der T -Verteilung gegen diejenige der standardisierten Normalverteilung ($N(0, 1)$).

• Pour $n \rightarrow \infty$ la fonction de distribution $f(x)$ ($z \rightsquigarrow x$) de la distribution de Student approche la fonction de la distribution normale standardisée ($N(0, 1)$).

$$f(x) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \Gamma(\frac{n}{2})} \cdot \frac{1}{(1 + \frac{x^2}{n})^{(n+1)/2}} \rightarrow \varphi(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$$

4.12.6 F-Verteilung von Fisher — Distribution de Fisher

Gegeben seien zwei voneinander unabhängige Stichproben mit den Umfängen n_1 und n_2 . Damit werden die Stichprobenfunktionen (Zufallsvariablen) $\bar{X}_1, \bar{X}_2, S_1^2, S_2^2$ gebildet: • *Soient données deux échantillons indépendants l'un de l'autre dont les tailles sont n_1 et n_2 . On les utilise pour former les fonctions d'échantillon (variables aléatoires) $\bar{X}_1, \bar{X}_2, S_1^2, S_2^2$:*

$$\begin{aligned}\bar{X}_1 &= \frac{1}{n_1} \sum_{i=1}^{n_1} X_{1i}, & S_1^2 &= \frac{1}{n_1 - 1} \sum_{i=1}^{n_1} (X_{1i} - \bar{X}_1)^2 \\ \bar{X}_2 &= \frac{1}{n_2} \sum_{i=1}^{n_2} X_{2i}, & S_2^2 &= \frac{1}{n_2 - 1} \sum_{i=1}^{n_2} (X_{2i} - \bar{X}_2)^2\end{aligned}$$

S_1^2 und S_2^2 sind unabhängig. • S_1^2 et S_2^2 sont indépendantes.

Seien die zugehörigen Grundgesamtheiten normalverteilt mit den Parametern $\mu_1, \mu_2, \sigma_1^2, \sigma_2^2$. Zudem setzen wir voraus:

• *Soient les données de bases liées distribuées de façon normale aux paramètres $\mu_1, \mu_2, \sigma_1^2, \sigma_2^2$. En plus nous supposons:*

Vor.: • Hyp.: $\sigma_1^2 = \sigma_2^2$

Nun bilden wir die folgende Stichprobenfunktion: • *Maintenant nous formons la fonction d'échantillon suivante:*

$$\mathcal{F} = \frac{S_1^2}{S_2^2} = \frac{\left(\frac{\chi_1^2 \sigma_1^2}{n_1 - 1}\right)}{\left(\frac{\chi_2^2 \sigma_2^2}{n_2 - 1}\right)} = \frac{\left(\frac{\chi_1^2}{n_1 - 1}\right)}{\left(\frac{\chi_2^2}{n_2 - 1}\right)} := \frac{\left(\frac{\chi_1^2}{m_1}\right)}{\left(\frac{\chi_2^2}{m_2}\right)}, \quad m_1 = n_1 - 1, \quad m_2 = n_2 - 1$$

Definition: • **Définition:** Die zu $\mathcal{F}$ gehörige Verteilung heisst **Fisher-Verteilung** oder **$\mathcal{F}$ -Verteilung** mit (m_1, m_2) Freiheitsgraden.

• *La distribution lié à $\mathcal{F}$ s'appelle **distribuiton de Fisher** ou **distribution $\mathcal{F}$** avec (m_1, m_2) degrés de liberté.*

Die Verteilungsfunktion und die Dichte berechnen sich wie folgt: • *On calcule la fonction de distribution et la densité comme il suit:*

Satz: • **Théorème:** Vor.: • Hyp.: $\mathcal{F}$ -Verteilung • *Distribution $\mathcal{F}$*

Beh.: • Thè.:

$$1. \quad F(x) = P(\mathcal{F} \leq x) = \frac{\Gamma\left(\frac{m_1 + m_2}{2}\right)}{\Gamma\left(\frac{m_1}{2}\right) \Gamma\left(\frac{m_2}{2}\right)} \cdot m_1^{\frac{m_1}{2}} \cdot m_2^{\frac{m_2}{2}} \cdot \int_0^x \frac{t^{\frac{m_1}{2}-2}}{(m_1 t + m_2)^{\frac{m_1 + m_2}{2}}} dt$$

$$2. \quad f(x) = \begin{cases} 0 & x \leq 0 \\ f(x) = E_{m_1, m_2} x^{\frac{m_1}{2}-1} (m_2 + m_1 \cdot x)^{-\frac{m_1 + m_2}{2}} & x > 0 \end{cases}$$

(f aus F durch ableiten!) • (f de F par calcul de la dérivée!)

4.13 Anhang I: Einige Beweise — Annexe I: Certaines preuves

4.13.1 Formel zur Gammafunktion — Formule pour la fonction gamma

Im nächsten Abschnitt brauchen wir folgenden Satz:

• Dans la section suivante nous avons besoin du théorème suivant:

Satz: • **Théorème:** **Vor.:** • **Hyp.:** $a, b > 0$

Beh.: • **Thè.:**

$$\frac{\Gamma(a) \cdot \Gamma(b)}{\Gamma(a+b)} = \int_0^1 u^{a-1} \cdot (1-u)^{b-1} du$$

Beweis: • **Preuve:**

Es ist: • Nous avons défini: $\Gamma(\alpha) = \int_0^\infty e^{-t} t^{\alpha-1} dt$

$$\leadsto \Gamma(a) \cdot \Gamma(b) = \left(\int_0^\infty e^{-t} t^{a-1} dt \right) \cdot \left(\int_0^\infty e^{-v} v^{b-1} dv \right) = \int_0^\infty \int_0^\infty e^{-(t+v)} t^{a-1} v^{b-1} dt dv = \int_0^\infty \int_0^\infty h(t, v) dt dv$$

$$\text{Sei } \bullet \text{ Soit } t = r \cdot u, \quad v = r(1-u) \Rightarrow \frac{\partial(t, v)}{\partial(r, u)} = \begin{vmatrix} \frac{\partial t}{\partial r} & \frac{\partial t}{\partial u} \\ \frac{\partial v}{\partial r} & \frac{\partial v}{\partial u} \end{vmatrix} = \begin{vmatrix} u & r \\ 1-u & -r \end{vmatrix} = -ur + r - (-ur) = -r$$

$$t = r \cdot u, \quad v = r(1-u), \quad t, r \in [0, \infty) \Rightarrow r \in [0, \infty), \quad u \in [0, 1] \quad (\text{Durch Betrachtung der Grenzfälle})$$

• (Par considération des cas limites)

$$\Rightarrow \Gamma(a) \cdot \Gamma(b) = \int_0^\infty \int_0^\infty h(t(r, u), v(r, u)) \cdot \left| \frac{\partial(t, v)}{\partial(r, u)} \right| dr du = \int_0^1 \int_0^\infty e^{-(r)} (ru)^{a-1} (r(1-u))^{b-1} r dr du$$

$$= \int_0^1 \int_0^\infty u^{a-1} (1-u)^{b-1} r dr du = \left(\int_0^\infty e^{-(r)} (r)^{a+b-1} r dr \right) \left(\int_0^1 e^{-(r)} (r)^{a+b-1} u^{a-1} (1-u)^{b-1} du \right) =$$

$$\Gamma(a+b) \Rightarrow \frac{\Gamma(a) \cdot \Gamma(b)}{\Gamma(a+b)} = \int_0^1 e^{-(r)} (r)^{a+b-1} u^{a-1} (1-u)^{b-1} du$$

☺

4.13.2 Dichte der Chi-Quadrat-Verteilung — Densité de la distribution Khi-deux

$X_i : \in N(0, 1) \leadsto$ normalverteilt • distribution normale, $\mu = 0, \sigma^2 = 1$. Es gilt: • Il vaut:

$$1. \quad x < 0 \Rightarrow P(\chi^2 = X_1^2 + \dots + X_n^2 = x < 0) = 0 \Rightarrow (x < 0 \Rightarrow f(x) = 0)$$

$$2. \quad (0 \leq x_i^2 \leq x \Leftrightarrow -\sqrt{x} \leq X_i \leq \sqrt{x}) \Rightarrow P(X_i^2 \leq x) = P(0 \leq X_i^2 \leq x) = P(-\sqrt{x} \leq X_i \leq \sqrt{x})$$

$$3. \quad X_i : \in N(0, 1) \Rightarrow$$

$$\forall_i F_1(x) = F_i(x) = P(X_i^2 \leq x) = P(-\sqrt{x} \leq X_i \leq \sqrt{x}) = 2 \cdot P(0 \leq X_i \leq \sqrt{x}) = \frac{2}{\sqrt{2\pi}} \int_0^{\sqrt{x}} e^{-\frac{u^2}{2}} du$$

$$\text{Sei } \bullet \text{ Soit } v = u^2, \quad u = v^{\frac{1}{2}} \Rightarrow \frac{du}{dv} = \frac{1}{2} v^{-\frac{1}{2}}, \quad du = dv \frac{1}{2} v^{-\frac{1}{2}} = dv \frac{1}{2\sqrt{v}}$$

$$\Rightarrow F_1(x) = \frac{2}{\sqrt{2\pi}} \frac{1}{2} \int_0^{\sqrt{v}=u=\sqrt{x}} e^{-\frac{v}{2}} \cdot \frac{1}{\sqrt{v}} dv = \frac{1}{\sqrt{2\pi}} \int_0^{\sqrt{v}=u=\sqrt{x}} \frac{e^{-\frac{v}{2}}}{\sqrt{v}} dv$$

$$\Rightarrow f_1(x) = \frac{dF_1(x)}{dx} = \frac{1}{\sqrt{2\pi}} \frac{e^{-\frac{x}{2}}}{\sqrt{x}}, \quad x > 0$$

$$4. \quad n = 1 \Rightarrow \chi^2 = X_1^2, \leadsto f(x) = K_1 x^{\frac{1-2}{2}} e^{-\frac{x}{2}} = \frac{1}{2^{\frac{1}{2}} \Gamma(\frac{1}{2})} \cdot x^{\frac{1-2}{2}} e^{-\frac{x}{2}} \stackrel{|\Gamma(\frac{1}{2})|=\sqrt{2}}{=} \frac{1}{\sqrt{2\pi}} \frac{1}{\sqrt{x}} \cdot e^{-\frac{x}{2}} \quad \text{☺}$$

Damit ist die Formel für $n = m = 1$ bewiesen. Wir können somit einen Induktionsbeweis versuchen.

• *On a donc prouvée la formule pour $n = m = 1$. Maintenant nous pouvons essayer de prouver la formule générale par induction.*

Sei • *Soit* $\chi^2 = X_1^2 + \dots + X_n^2 := \chi_n^2, \quad f(x) := f_n(x)$

Sei die Formel richtig für $m=n-1$. • *La formule soit juste pour $m = n - 1$.*

$$\leadsto f_{n-1}(x) = K_{n-1} x^{\frac{n-1-2}{2}} e^{-\frac{x}{2}} = \frac{1}{2^{\frac{n-1}{2}} \Gamma(\frac{n-1}{2})} \cdot x^{\frac{n-1-2}{2}} e^{-\frac{x}{2}} = \frac{1}{2^{\frac{n-1}{2}} \Gamma(\frac{n-1}{2})} \cdot x^{\frac{n-3}{2}} e^{-\frac{x}{2}}$$

Es gilt: • *Il vaut:* $\chi_n^2 = \chi_{n-1}^2 + X_n^2 \quad \chi_{n-1} \leftrightarrow f_{n-1}(x) \wedge \chi_n \leftrightarrow f_n(x) = f_1(x)$

Nach Voraussetzung sind die X_i unabhängig und daher sind auch χ_{n-1} und X_n^2 unabhängig. Daher können wir die folgende Formel für die Dichte von Seite 141 benutzen:

• *D'après l'hypothèse les X_i sont indépendantes. Par conséquent, χ_{n-1} et X_n^2 sont aussi indépendantes. Nous pouvons donc utiliser la formule suivante pour la densité de la page 141:*

$$f_n(x) = \int_{-\infty}^{\infty} f_n(x-y) \cdot f_{n-1}(y) dy = \int_0^x f_1(x-y) \cdot f_{n-1}(y) dy$$

($\int_{-\infty}^{\infty} \dots = \int_0^x$ wegen • *à cause de $y < 0 \Rightarrow f_{n-1}(y) = 0, \quad x > y \Rightarrow x-y < 0 \Rightarrow f_1(x-y) = 0$*)

$$\begin{aligned} \leadsto f_n(x) &= \int_0^x f_1(x-y) \cdot f_{n-1}(y) dy = \frac{1}{\sqrt{2\pi}} \cdot \frac{1}{2^{\frac{n-1}{2}} \Gamma(\frac{n-1}{2})} \cdot \int_0^x \frac{e^{-\frac{x-y}{2}}}{\sqrt{x-y}} \cdot y^{\frac{n-3}{2}} \cdot e^{-\frac{y}{2}} dy \\ &= \frac{1}{2^{\frac{n}{2}} \cdot \sqrt{\pi} \cdot \Gamma(\frac{n-1}{2})} \cdot e^{-\frac{x}{2}} \cdot \int_0^x (x-y)^{-\frac{1}{2}} \cdot y^{\frac{n-3}{2}} dy := \frac{1}{2^{\frac{n}{2}} \cdot \sqrt{\pi} \cdot \Gamma(\frac{n-1}{2})} \cdot e^{-\frac{x}{2}} \cdot I, \quad I = \int_0^x (x-y)^{-\frac{1}{2}} \cdot y^{\frac{n-3}{2}} dy \end{aligned}$$

Sei • *Soit* $y := ux, \quad y \in [0, x] \Rightarrow u = \frac{y}{x} \in [0, 1], \quad x-y = x-ux = x(1-u), \quad \frac{dy}{du} = x, \quad dy = du x$

$$\leadsto I = \int_0^x (x(1-u))^{-\frac{1}{2}} \cdot (ux)^{\frac{n-3}{2}} x du = (x^{-\frac{1}{2}} x^{\frac{n-3}{2}} x) \int_0^x (1-u)^{-\frac{1}{2}} \cdot (u)^{\frac{n-3}{2}} du = x^{\frac{n-2}{2}} \int_0^x (1-u)^{-\frac{1}{2}} \cdot u^{\frac{n-3}{2}} du$$

Satz über die Gammafunktion von Seite 149 • *Théorème sur la fonction gamma de la page 149*

$$\leadsto I = x^{\frac{n-2}{2}} \int_0^x (1-u)^{\frac{1}{2}-1} \cdot u^{\frac{n-1}{2}-1} du = x^{\frac{n-2}{2}} \frac{\Gamma(\frac{n-1}{2})}{\Gamma(\frac{1}{2})} \Gamma(\frac{n-1}{2} + \frac{1}{2}), \quad a = \frac{1}{2} > 0, \quad b = \frac{n-1}{2} > 0, \quad (n > 2)$$

$$\Gamma(\frac{1}{2}) = \sqrt{\pi} \Rightarrow f_n(x) = \frac{1}{2^{\frac{n}{2}} \cdot \sqrt{\pi} \cdot \Gamma(\frac{n-1}{2})} \cdot e^{-\frac{x}{2}} \cdot x^{\frac{n-2}{2}} \frac{\Gamma(\frac{n-1}{2}) \cdot \sqrt{\pi}}{\Gamma(\frac{n-1}{2} + \frac{1}{2})} = \frac{1}{2^{\frac{n}{2}}} \cdot e^{-\frac{x}{2}} \cdot x^{\frac{n-2}{2}} \frac{1}{\Gamma(\frac{n}{2})} = f(x) \quad \text{☺}$$

Das ist die behauptete Formel. • *C'est la formule qu'on a dû démontrer.*

4.13.3 Dichte der Student-Verteilung — Densité de la distribution de Student

Verteilungen nach Voraussetzung: • *Distributions selon hypothèse:*

$$X \rightsquigarrow N(0, 1) \rightsquigarrow X \leftrightarrow f(x) := f_1(x) = \frac{1}{\sqrt{2\pi}} e^{-\frac{x^2}{2}}$$

$$Y \rightsquigarrow Y = \chi_n^2 \rightsquigarrow X \leftrightarrow f(y) := f_2(y) = \frac{1}{2^{\frac{n}{2}}} \cdot e^{-\frac{y}{2}} \cdot y^{\frac{n-2}{2}} \frac{1}{\Gamma(\frac{n}{2})}, \quad y > 0 \quad (y \leq 0 \Rightarrow f_2(y) = 0)$$

X, Y unabhängig • *indépendantes* $\rightsquigarrow f(x, y) = f_1(x) f_2(y)$

$$F(z) := P(T \leq z) = P\left(\frac{X}{Y/n} \leq z\right) = P(X \leq z \cdot Y/n) \rightsquigarrow F(z) = \int \int_{x \leq z \cdot y/n} f(x, y) dx dy =$$

$$= \int \int_{x \leq z \cdot \sqrt{y/n}, \quad y > 0} f_1(x) f_2(y) dx dy = \frac{1}{\sqrt{2\pi}} \cdot \frac{1}{2^{\frac{n}{2}}} \cdot \frac{1}{\Gamma(\frac{n}{2})} \int \int_{x \leq z \cdot \sqrt{y/n}, \quad y > 0} e^{-\frac{x^2}{2}} \cdot e^{-\frac{y}{2}} \cdot y^{\frac{n-2}{2}} dx dy =$$

$$= C_n \int \int_{x \leq z \cdot \sqrt{y/n}, \quad y > 0} e^{-\frac{x^2}{2}} \cdot e^{-\frac{y}{2}} \cdot y^{\frac{n-2}{2}} dx dy = C_n \int_0^\infty \left(\int_{-\infty}^{x=z \cdot \sqrt{y/n}} e^{-\frac{x^2}{2}} \cdot e^{-\frac{y}{2}} \cdot y^{\frac{n-2}{2}} dx \right) dy,$$

$$C_n = \frac{1}{\sqrt{2\pi} \cdot 2^{\frac{n}{2}} \cdot \Gamma(\frac{n}{2})} \quad \text{Sei } \bullet \text{ Soit } x := u \sqrt{\frac{y}{n}} \Rightarrow dx = du \sqrt{\frac{y}{n}}, \quad H := 1 + \frac{u^2}{n}, \quad \frac{H y}{2} = \frac{y}{2} + \frac{x^2}{n}$$

$$\rightsquigarrow F(z) = \frac{C_n}{\sqrt{n}} \int_0^\infty \left(\int_{-\infty}^{u \sqrt{\frac{y}{n}} = x = z \cdot \sqrt{y/n}} e^{-\frac{(u \sqrt{\frac{y}{n}})^2}{2}} \cdot e^{-\frac{y}{2}} \cdot y^{\frac{n-2}{2}} du \right) \sqrt{y} dy = \frac{C_n}{\sqrt{n}} \int_0^\infty \left(\int_{-\infty}^{u=z} e^{-\frac{\overbrace{(x)^2}^{H y}}{2}} \cdot y^{\frac{n-1}{2}} dy \right) du$$

$$= \frac{C_n}{\sqrt{n}} \int_0^\infty \left(\int_{-\infty}^{u=z} e^{-\frac{H y}{2}} \cdot y^{\frac{n-1}{2}} dy \right) du \quad \text{Sei } \bullet \text{ Soit } H y := 2v, \quad y = \frac{2v}{H}, \quad dy = dv \frac{2}{H} \rightsquigarrow$$

$$\rightsquigarrow F(z) = \frac{C_n}{\sqrt{n}} \int_{v=0}^{v=\infty} \left(\int_{u=-\infty}^{u=z} e^{-\frac{2v}{H}} \cdot \left(\frac{2v}{H}\right)^{\frac{n-1}{2}} \cdot \frac{2}{H} dv \right) du = \frac{C_n}{\sqrt{n}} \int_{v=0}^{v=\infty} \left(\frac{1}{H}\right)^{\frac{n-1}{2}} \cdot \frac{2}{H} \left(\int_{u=-\infty}^{u=z} e^{-\frac{2v}{H}} \cdot (2v)^{\frac{n-1}{2}} dv \right) du$$

$$= \frac{C_n}{\sqrt{n}} \cdot 2^{\frac{n+1}{2}} \cdot \left(\int_{v=0}^{v=\infty} \left(\frac{1}{H}\right)^{\frac{n+1}{2}} du \right) \cdot \left(\int_{u=-\infty}^{u=z} e^{-\frac{2v}{H}} \cdot v^{\frac{n-1}{2}} dv \right) = \frac{C_n}{\sqrt{n}} \cdot 2^{\frac{n+1}{2}} \cdot \left(\int_{u=-\infty}^{u=z} \left(\frac{1}{H}\right)^{\frac{n+1}{2}} du \right) \cdot \Gamma\left(\frac{n+1}{2}\right)$$

$$= \frac{1}{\sqrt{2\pi} \cdot 2^{\frac{n}{2}} \cdot \Gamma(\frac{n}{2})} \cdot \frac{1}{\sqrt{n}} \cdot 2^{\frac{n+1}{2}} \cdot \left(\int_{u=-\infty}^{u=z} \left(\frac{1}{1 + \frac{u^2}{n}}\right)^{\frac{n+1}{2}} du \right) \cdot \Gamma\left(\frac{n+1}{2}\right) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \cdot \Gamma(\frac{n}{2})} \cdot \left(\int_{u=-\infty}^{u=z} \left(\frac{1}{1 + \frac{u^2}{n}}\right)^{\frac{n+1}{2}} du \right)$$

$$\rightsquigarrow F(z) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \cdot \Gamma(\frac{n}{2})} \cdot \left(\int_{u=-\infty}^{u=z} \left(\frac{1}{1 + \frac{u^2}{n}}\right)^{\frac{n+1}{2}} du \right), \quad f(z) = \frac{dF(z)}{dz} = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n\pi} \cdot \Gamma(\frac{n}{2})} \cdot \left(\frac{1}{1 + \frac{z^2}{n}}\right)^{\frac{n+1}{2}} \quad \text{☺}$$

4.13.4 Beweis Tschebyscheffsche Ungleichung — Preuve d'inéquation de Tschebyscheff

Sei Y eine beliebige Zufallsgrösse, die nicht normalverteilt sein muss. $E(Y) = \mu$ sei hier der Erwartungswert, $\text{Var}(Y) = \sigma^2$ sei hier die Varianz und ε sei eine beliebige positive Zahl. Dann gilt:

• *Soit Y une variable aléatoire quelconque, qui ne doit pas être distribuée de façon normale. Soit ici $E(Y) = \mu$ la valeur d'espérance, $\text{Var}(Y) = \sigma^2$ soit ici la variance et ε soit un nombre positif et quelconque. Alors il vaut:*

Formel: • **Formule:** **Tschebyscheff**

$$P(|Y - \mu| \geq \varepsilon) \leq \frac{\text{Var}(Y)}{\varepsilon^2}$$

Bemerkung: • **Remarque:** $P(|Y - \mu| \geq \varepsilon) \leq \frac{\sigma^2}{\varepsilon^2} \Leftrightarrow P(|Y - \mu| < \varepsilon) > 1 - \frac{\sigma^2}{\varepsilon^2}$

Definition: • **Définition:** **Nichtnegative Zufallsvariable X :** Nimmt mit Wahrscheinlichkeit 1 nur Werte aus $\mathbb{R}_0^+$ an.
• Variable aléatoire non-négative X : N'atteint que des valeurs $\mathbb{R}_0^+$ avec la probabilité 1.

Wir beweisen erst ein Lemma: • *Nous prouvons d'abord un lemme:*

Lemma: • **Lemme:** **Vor.:** • **Hyp.:** X nichtnegativ • *non-négative*, $\delta > 0$

Beh.: • **Thè.:** $\frac{E(X)}{\delta} \leq P(X \geq \delta)$

Beweis: • **Preuve:** (Lemma) • (*Lemme*)

1. Diskreter Fall: • *Cas discret:*

$$E(X) = \sum_{k \in \{k\}} x_k p_k \geq \sum_{k \in \{k \mid x_k \geq \delta\} := H} x_k p_k \geq \sum_{k \in H} x_{\min} p_k \geq \sum_{k \in H} \delta p_k \geq \delta \sum_{k \in H} p_k = \delta P(X \geq \delta)$$

2. Stetiger Fall: • *Cas continu:*

$$\begin{aligned} E(X) &= \int_{-\infty}^{\infty} x f(x) dx \geq \int_0^{\infty} x f(x) dx \geq \int_{\delta}^{\infty} x f(x) dx \geq \int_{\delta}^{\infty} x_{\min} f(x) dx \geq \int_{\delta}^{\infty} \delta f(x) dx \\ &= \delta \int_{\delta}^{\infty} f(x) dx = \delta P(X \geq \delta) \end{aligned}$$

Beweis: • **Preuve:** (Satz) • (*Théorème*)

Sei • *Soit* $\delta = \varepsilon^2$, $X = |Y - \mu|^2 \rightsquigarrow$ nichtnegativ • *non-négative*

$$\begin{aligned} \sigma^2 &= E(|Y - \mu|^2) = E((Y - E(Y))^2) = E(X) \geq \delta P(X \geq \delta) = \varepsilon^2 P(|Y - E(Y)|^2 \geq \varepsilon^2) \\ &= \varepsilon^2 P(|Y - E(Y)| \geq \varepsilon) \Rightarrow P(|Y - E(Y)| \geq \varepsilon) \leq \frac{\sigma^2}{\varepsilon^2} \end{aligned}$$

4.14 Anhang II: Ergänzungen — Annexe II: Suppléments

4.14.1 Quadratsumme contra Betragssumme — La somme des carrés contre somme de valeur absolue

1. **Frage:** Für welches c ist die Quadratsumme $\sum_i (x_i - c)^2$ minimal?

• **Question:** Pour quel c la somme des carrés $\sum_i (x_i - c)^2$ devient-elle minimale?

$$\begin{aligned} \leadsto \frac{d}{dc} \sum_{i=1}^n (x_i - c)^2 &= \sum_{i=1}^n (-1) \cdot 2(x_i - c) = -2 \sum_{i=1}^n (x_i - c) = 0 \Rightarrow \sum_{i=1}^n x_i = \sum_{i=1}^n c = n \cdot c \\ \Rightarrow c &= \frac{1}{n} \cdot \sum_{i=1}^n x_i = \bar{x} \Rightarrow c = \bar{x} \leadsto c = \text{Mittelwert} \bullet c = \text{moyenne} \end{aligned}$$

2. **Frage:** Für welches c ist die Betragssumme $\sum_i |x_i - c|$ minimal?

• **Question:** Pour quel c la somme des valeurs absolue $\sum_i |x_i - c|$ devient-elle minimale?

$$\begin{aligned} \leadsto \frac{d}{dc} \sum_{i=1}^n |x_i - c| &= \frac{d}{dc} \sum_{i=1}^n \operatorname{sgn}(x_i - c) \cdot (x_i - c) = \sum_{i=1}^n \operatorname{sgn}(x_i - c) = \sum_{i=1}^n (-1) \cdot \operatorname{sgn}(x_i - c) \\ &= -(\underbrace{-1 - 1 - \dots - 1}_{x_i < c}) + (\underbrace{+1 + 1 + \dots + 1}_{x_i > c}) = 0 \\ \Rightarrow -(\underbrace{-1 - 1 - \dots - 1}_{x_i < c}) &= (\underbrace{+1 + 1 + \dots + 1}_{x_i > c}) \\ \leadsto (\text{Anzahl „Werte } < c\text{“}) &= (\text{Anzahl „Werte } > c\text{“}) \leadsto c = \text{Median} \\ \bullet (\text{Nombre de „valeurs } < c\text{“}) &= (\text{nombre de „valeurs } > c\text{“}) \leadsto c = \text{médian} \end{aligned}$$

Konsequenz: • **Conséquence:**

$\sum_i (x_i - c)^2$ wird minimal für den Mittelwert $c = \bar{x}$, $\sum_{i=1}^n |x_i - c|$ wird minimal für den Median $c = \tilde{x}$.

• $\sum_i (x_i - c)^2$ devient minimal pour la moyenne $c = \bar{x}$, $\sum_{i=1}^n |x_i - c|$ est minimale pour le médian $c = \tilde{x}$.

4.14.2 Verteilungstreue u.s.w. — Conformité du type de répartition etc.

Auf Seite 141 haben wir folgenden Satz gesehen:

• A la page 141 on a vu le théorème suivant:

Satz: • **Théorème:**

Summe unabhängiger normalverteilter Variablen

• **Somme de variables indépendantes et distribuées de façon normale**

Vor.: • **Hyp.:**

Seien $X_1, X_2, \dots, X_n$ unabhängig, normalverteilt mit

• Soient $X_1, X_2, \dots, X_n$ indépendantes, distribuées de façon normale avec

Mittelwerte • Moyennes $\mu_1, \mu_2, \dots, \mu_n$

Varianzen • Variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$

Beh.: • **Thè.:**

$\sum_{i=1}^n X_i$ normalverteilt • distribuées de façon normale

Mittelwert • Moyenne $\mu = \sum_{i=1}^n \mu_i$

Varianz • Variance $\sigma^2 = \sum_{i=1}^n \sigma_i^2$

Summen unabhängiger normalverteilter Variablen sind somit wieder normalverteilt. Dieses Verhalten tritt aber auch bei anderen Verteilungsarten zutage:

- Les sommes de variables distribuées de façon normale sont de nouveau distribuées de façon normale. On observe ce comportement aussi aux autres types de répartition:

Satz:	• Théorème:	<u>Vor.:</u>	• Hyp.:	X, Y unabhängig • indépendantes $X : \in Bi(n_1, p), Y : \in Bi(n_2, p)$
		<u>Beh.:</u>	• Thè.:	$X + Y : \in Bi(n_1 + n_2, p)$

Zum Beweis in einem Spezialfall: • Quant à la preuve dans un cas spécial:

Sei X = Anzahl des Eintretens des Ereignisses A mit $p(A) = p$ bei n_1 -maligem Wiederholen des zugehörigen Versuches und Y = Anzahl des Eintretens des Ereignisses A mit $p(A) = p$ bei n_2 -maligem Wiederholen des zugehörigen Versuches. Offensichtlich gehört dann $X + Y$ zum selben Versuch mit $p(A) = p$: $X + Y$ = Anzahl des Eintretens des Ereignisses A bei $n_1 + n_2$ -maligem Wiederholen des zugehörigen Versuches. Die Verteilung ist immer eine Binomialverteilung. (Zum Beweis vgl auch Lit. A4)

- Soit X = nombre de réalisations de l'événement A avec $p(A) = p$ si l'expérience affilée est répétée n_1 fois et Y = nombre de réalisations de l'événement A avec $p(A) = p$ si l'expérience affilée est répétée n_2 fois. Évidemment $X + Y$ est liée à la même expérience avec $p(A) = p$: $X + Y$ = nombre de réalisations de l'événement A si l'expérience affilée est répétée $n_1 + n_2$ fois. La répartition est toujours une répartition de Bernoulli (binomiale). (Quant à la preuve voir aussi lit. A4)

Da die Binomialverteilung sich bei $np = \lambda = \text{const.}$ einer Poissonverteilung annähert, lässt sich der Satz auch auf die Poissonverteilung übertragen:

- Comme la répartition binomiale sous la condition $np = \lambda = \text{const.}$ s'approche à une répartition de Poisson, on peut adapter le théorème aussi à une répartition de Poisson:

Satz:	• Théorème:	<u>Vor.:</u>	• Hyp.:	X, Y unabhängig • indépendantes $X : \in Po(\lambda_1), Y : \in Bi(\lambda_2)$
		<u>Beh.:</u>	• Thè.:	$X + Y : \in Bi(\lambda_1 + \lambda_2)$

Für die Beweise der folgenden Sätze sei auf die Literatur verwiesen (falls nicht schon erledigt):

- Quant aux preuves des théorèmes suivants le lecteur est renvoyé à la littérature (si ce n'est pas encore terminé):

Satz:	• Théorème:	<u>Vor.:</u>	• Hyp.:	X, Y unabhängig • indépendantes X, Y χ^2 -verteilt • réparties de façon χ^2 $X : \in \chi_{m_1}^2, Y : \in \chi_{m_2}^2$
		<u>Beh.:</u>	• Thè.:	$X + Y : \chi_{m_1+m_2}^2$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$X_1, X_2, \dots, X_n$ unabhängig • *indépendantes*
 X_i $N(0, 1)$ -verteilt • *réparties de façon $N(0, 1)$*

Beh.: • **Thè.:**

$$Z = X_1^2 + X_2^2 + \dots + X_n^2 : \in \chi_n^2$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$X_1, X_2, \dots, X_n$ unabhängig • *indépendantes*
 X_i $N(\mu, \sigma^2)$ -verteilt • *réparties de façon $N(\mu, \sigma^2)$*

$$\bar{X} = \frac{1}{n} \sum_{i_1}^n X_i$$

Beh.: • **Thè.:**

$$\frac{1}{\sigma^2} \sum_{i_1}^n (X_i - \bar{X})^2 : \in \chi_{n-1}^2$$

Bezüglich der t -Verteilung und der χ^2 -Verteilung können wir noch rekapitulierend zusammenfassen oder folgern:

• *Par rapport à la répartition de Student et de khi-deux, nous pouvons déduire en récapitulant ou comme résumé:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

X, Y unabhängig • *indépendantes*
 $X : \in N(0, 1), Y : \in \chi_n^2$
 $Z = \sqrt{n} \cdot \frac{X}{\sqrt{Y}}$

Beh.: • **Thè.:**

$$Z : \in t_n$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$X_1, X_2, \dots, X_n$ unabhängig • *indépendantes*
 $X_i : N(\mu, \sigma^2)$ -verteilt • *réparties de façon $N(\mu, \sigma^2)$*
 $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i, Y = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2, Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}}$

Beh.: • **Thè.:**

$$\bar{X} : \in N\left(\mu, \frac{\sigma}{\sqrt{n}}\right), Z : \in N(0, 1), Y : \in \chi_{n-1}^2$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\begin{aligned}
 &X_1, X_2, \dots, X_n \text{ unabhängig} \bullet \text{ indépendantes} \\
 &X_i: N(\mu, \sigma^2)\text{-verteilt} \bullet \text{ réparties de façon } N(\mu, \sigma^2) \\
 &H = \sqrt{n} \cdot \frac{\bar{X} - \mu}{\sqrt{\frac{\sum_{i=1}^n (X_i - \bar{X})^2}{n-1}}}
 \end{aligned}$$

Beh.: • **Thè.:**

$$H : \in t_{n-1}$$

4.14.3 Zentraler Grenzwertsatz — Théorème limite central

Auf Seite 120 haben wir den Grenzwertsatz von Moivre/ Laplace kennengelernt. Bei n Zufallsvariablen sieht er wie folgt aus:

• A la page 120 on a vu le théorème de Moivre/ Laplace. Avec n variables aléatoires il a la forme suivante:

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\begin{aligned}
 &\text{Sei} \bullet \text{ Soit } H_n : \in \text{Bi}(n, p), \quad p \in (0, 1), \quad n \in \mathbb{N} \\
 &Y_n := \frac{H_n - E(H_n)}{\sqrt{\text{Var}(H_n)}} = \frac{H_n - E(H_n)}{\sqrt{np(1-p)}} \\
 &\leadsto Y_n \text{ standardisiert} \bullet \text{ standardisée} \\
 &F_{Y_n} \text{ Verteilungsfunktion} \bullet \text{ fonction de distribution}
 \end{aligned}$$

Beh.: • **Thè.:**

$$\lim_{n \rightarrow \infty} F_{Y_n} = \Phi(x) = \frac{1}{2\pi} \int_{-\infty}^x e^{-\frac{t^2}{2}} dt, \quad x \in \mathbb{R}$$

Bemerkung: • **Remarque:**

Der Satz hat eine Bedeutung, wenn es bei grossen n um approximative Berechnungen von Wahrscheinlichkeiten im Zusammenhang mit Binomialverteilungen geht. Hier ist dann $\mu = np$, $\sigma = \sqrt{np(1-p)}$. Dann gilt:

• Le théorème est important si pour des n grands, il s'agit de calculer approximativement des probabilités dans le contexte de distributions binomiales. Il est donc $\mu = np$, $\sigma = \sqrt{np(1-p)}$. Alors il vaut:

Folgerung: • **Conclusion:**

$$P(a \leq X \leq b) = P\left(\frac{a - np}{\sqrt{np(1-p)}} \leq \frac{X - np}{\sqrt{np(1-p)}} \leq \frac{b - np}{\sqrt{np(1-p)}}\right) \approx \Phi\left(\frac{b - np}{\sqrt{np(1-p)}}\right) - \Phi\left(\frac{a - np}{\sqrt{np(1-p)}}\right)$$

Φ ist leicht bestimmbar (Rechner, Tabellen,...). Gute Näherungen erreicht man schon, wenn die **Laplace-Bedingung** $np(1-p) > 9$ erfüllt ist.

• Il est facile d'évaluer Φ (ordinateur, tableaux,...). L'approximation est déjà bonne si la condition $np(1-p) > 9$ est satisfaite **condition de Laplace**.

Wir schreiben $A \in R_i$, wenn das Ereignis A bei der i -ten Wiederholung eines Versuches eintritt. Sonst ist $A \notin R_i$. Wir wählen nun H_n speziell:

- *Nous écrivons $A \in R_i$, si l'événement A est réalisé lors de la i -ième répétition d'une expérience. Autrement nous écrivons $A \notin R_i$. Soit H_n la variable comme il suit:*

$$H_n(A) = \sum_{i=1}^n X_i, \quad X_i = \begin{cases} 1 & A \in R_i \\ 0 & A \notin R_i \end{cases}$$

Man sieht sofort, dass jedes X_i für sich binomialverteilt ist mit den Parametern $n = 1$ und p , d.h. mit dem Erwartungswert $\mu = np = p$ und der Varianz $\sigma^2 = np(1-p) = p(1-p)$. Wie wir schon wissen, ist die Variable $\sum_{i=1}^n$ als Summe binomialverteilter Zufallsvariablen wieder binomialverteilt.

- *On voit tout de suite que chaque X_i pour lui-même est distribué de façon binomiale avec les paramètres $n = 1$ et p , ç.v.d. avec la valeur d'espérance $\mu = np = p$ et la variance $\sigma^2 = np(1-p) = p(1-p)$. Comme nous savons, la variable $\sum_{i=1}^n$ est distribuée de façon binomiale comme somme de variables distribuées de façon binomiale.*

Sei • Soit
$$Y_n = \frac{\sum_{i=1}^n X_i - n\mu}{\sigma\sqrt{n}} = \frac{\frac{1}{n} \sum_{i=1}^n X_i - \mu}{\sqrt{\sigma}/\sqrt{n}} \rightsquigarrow F_{Y_n} \rightarrow \Phi(n), \quad n \rightarrow \infty$$

Dieses Verhalten ist auch beweisbar, wenn $\langle X_n \rangle$ eine beliebige Folge unabhängiger und identisch verteilter Zufallsgrößen ist mit gemeinsamem Erwartungswert μ und gemeinsamer Varianz $\sigma^2 \in (0, k)$, $k = \text{const.} \in \mathbb{R}^+$.

- *Ce comportement est aussi prouvable si $\langle X_n \rangle$ est une suite quelconque de variables aléatoires indépendantes et identiquement distribuées avec une valeur d'espérance μ commune et une variance σ^2 commune, $\sigma^2 \in (0, k)$, $k = \text{const.} \in \mathbb{R}^+$.*

Satz: • **Théorème:** **Zentraler Grenzwertsatz**
• **Théorème limite central**

Vor.: • Hyp.:

Wie beschrieben • *Comme décrit*

Beh.: • Thè.:

$$\lim_{n \rightarrow \infty} F_{Y_n} = \Phi(x) = \frac{1}{2\pi} \int_{-\infty}^x e^{-\frac{t^2}{2}} dt, \quad x \in \mathbb{R}$$

Beweise finden sich in der Literatur, z.B. in Bibl. A3.

- *On trouve des preuves dans la littérature, p.ex. dans Bibl. A3.*

4.14.4 Lineare Transformationen — Transformations linéaires

Auf Seite 95 haben wir den Satz erkannt: • *A la page 95 nous avons vu le théorème:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$g(X) = a h(X) + b u(X)$$

Beh.: • **Thè.:**

$$E(g(X)) = a E(h(X)) + b E(u(X))$$

Und auf Seite 95 sahen wir: • *E à la page 95 nous avons vu:*

Satz: • **Théorème:** $E(a_n X^n + a_{n-1} X^{n-1} + \dots + a_1 X + a_0) = a_n E(X^n) + a_{n-1} E(X^{n-1}) + \dots + a_1 E(X) + a_0 = a_n \mu_n + a_{n-1} \mu_{n-1} + \dots + a_1 \mu_1 + a_0$

$\leadsto$ **Speziell:** • **Spécialement:** $E(X) = \mu, \quad X^* = c_1 X + c_2 = g(X)$

Korollar: • **Corollaire:** **Vor.:** • **Hyp.:**

$$X^* = c_1 X + c_2$$

Beh.: • **Thè.:**

$$\mu^* = c_1 \mu + c_2$$

Weiter können wir folgern: • *En outre nous déduisons:*

$$X^* - \mu^* = (c_1 X + c_2) - (c_1 \mu + c_2) = c_1 (X - \mu), \quad E((X^* - \mu^*)^k) = E(c_1^k (X - \mu)^k) = c_1^k E((X - \mu)^k)$$

$\leadsto$ **Speziell:** • **Spécialement:** $k = 2 : \sigma_{X^*}^2 = E((X^* - \mu^*)^2) = c_1^2 E((X - \mu)^2) = c_1^2 \sigma^2$

Korollar: • **Corollaire:** **Vor.:** • **Hyp.:**

$$X^* = c_1 X + c_2$$

Beh.: • **Thè.:**

$$\sigma_{X^*}^2 = c_1^2 \sigma^2$$

Kapitel 5

Mathematische Statistik — Statistique mathématique

5.1 Qualitätskontrolle — Contrôle de qualité

5.1.1 Allgemeines, SQC — Généralités, SQC

Qualitätskontrolle und speziell **statistische Qualitätskontrolle (SQC, statistical quality control)** umfasst Verfahren, die zur Qualitätssicherung, Qualitätsprüfung und Qualitätslenkung benutzt werden. Heute sind die Rahmenbedingungen oft Gegenstand von Normierungen (z.B. DIN 55350). Generell kann man zwischen qualitativen Kontrollen (z.B. gut, schlecht) und quantitativen Kontrollen (Messung von Variablen) unterscheiden.

- *Sous contrôle de qualité et spécialement contrôle de qualité statistique (SQC, statistical quality control) on entend des méthodes qui sont utilisées pour l'assurance de qualité, l'examen de la qualité et la gestion de qualité. Aujourd'hui, les conditions sont souvent objet de normalisations, par exemple D.I.N. 55350. Généralement on peut distinguer entre les contrôles qualitatifs ou contrôles par attributs (p.ex. bon, mauvais) et les contrôles quantitatifs ou contrôles par variables (mesurage de variables).*

Entsprechend der Anwendung unterscheidet man **zwei Hauptgebiete**:

- *Selon l'utilisation on distingue deux domaines principaux:*

1. Statistische Prozesskontrolle, Fertigungsüberwachung oder Qualitätsregulierung (**SPC, statistical process control**):

- *Contrôle statistique de processus, surveillance de fabrication ou réglage de qualité (SPC, statistical process control):*

Hier geht es um die Überwachung des Fertigungsprozesses bei der Herstellung von Massenprodukten, um während der Produktion Qualitätsstörungen entdecken und dann unmittelbar eingreifen und steuern zu können.

- *Ici, il s'agit de la surveillance d'un processus de fabrication pendant la production de produits de masse, pour découvrir des différences de qualité et pour pouvoir intervenir et conduire directement.*

2. Annahmekontrolle oder Annahmestichprobenprüfung (AC, acceptance sampling):

- *Contrôle de réception ou examen d'échantillon de réception (AC, acceptance sampling):*

Hier geht es um die Eingangs-, Zwischen- und Endkontrolle eines Betriebes zur Überprüfung der Erzeugnisse ohne direkten Einfluss auf die Produktion. Damit wird aber der Umfang von produziertem Ausschuss beziffert. Die Annahmekontrolle dient auch zur Rückweisung von eingehender Ware. Sie beeinflusst daher indirekt die Produktion.

- *Ici, il s'agit du contrôle d'entrée, d'un contrôle pendant la production et d'un contrôle final des produits dans une entreprise (ou usine) sans influence directe sur la production. Ainsi le montant*

de rebut produit est mesuré. Le contrôle initial sert aussi à refuser la marchandise arrivante. Elle n'influence par conséquent la production que de manière indirecte.

5.2 SQC1: Prozesskontrolle — SQC1: Contrôle de processus

5.2.1 Problemstellung — Problème

Bsp.: • Exemple:

Es sollen Sollwerte innerhalb gewisser Toleranzen eingehalten werden. Z.B. bei der Serienfertigung von Zahnrädern muss der Durchmesser der zentralen Bohrung einen bestimmten Sollwert μ_0 aufweisen.

• *Il faut observer des valeurs prescrites par rapport à certaines tolérances. P.ex. lors de la fabrication en série de roues dentées le diamètre du forage central doit montrer une valeur prescrite μ_0 .*

Aus Erfahrung weiss man aber, dass auch bei aller Sorgfalt stets Abweichungen auftreten. Gründe sind menschliches Versagen, Abnutzungserscheinungen bei den Maschinen und Apparaten u.s.w. sowie Materialprobleme. Daher gilt es, die Produktion permanent mittels SPC zu überwachen und gegebenenfalls sofort einzugreifen. Ziel ist es, den Ausschuss auf ein Minimum zu beschränken.

• *Par expérience on sait qu'il y a des écarts qui apparaissent malgré tout le soin. Les raisons sont la défaillance humaine, usures ou abrasions aux machines et appareils etc. ainsi que des problèmes de matériel. Par conséquent il est essentiel de surveiller la production de façon permanente au moyen de SPC et d'intervenir immédiatement s'il est nécessaire. Le but est de limiter le rebut à un minimum.*

5.2.2 Beispiel Mittelwert — Exemple moyenne

In der Praxis gibt es verschiedenste Problemstellungen, die wir hier nicht alle behandeln können. Wir greifen daher beispielhaft ein Problem heraus. Bei einer momentan mit einem Automaten produzierten Serie eines Artikels soll der Mittelwert μ_0 innerhalb eines nicht kritischen Toleranzbereiches konstant eingehalten werden.

• *Dans la pratique, il y a des problèmes différents que nous ne pouvons pas tous traiter ici. Par conséquent nous choisissons un problème comme exemplaire. Lors de la production automatique d'une série d'un article il faudrait constamment observer la moyenne μ_0 pour qu'elle reste dans une certaine marge de tolérance non critique.*

Methodisches Vorgehen: • Procédé méthodique (manière d'agir):

1. Annahme: • Hypothese:

- (a) Aus Erfahrung wissen wir und nehmen daher auch jetzt an: $\bar{X} = \mu$ genügt einer Normalverteilung mit bekannter Standardabweichung σ . Diese Annahme ist jedoch nur in einem gewissen Bereich vernünftig, da der Definitionsbereich der Normalverteilung $(-\infty, \infty) = \mathbb{R}$ ist, der Bereich, in dem Messwerte x_i und damit Mittelwerte $\bar{x}$ überhaupt feststellbar sind, ist $(0, M_{\max})$. $M_{\max}$ hängt von den praktischen Möglichkeiten ab. Negative Messwerte existieren nicht. μ_0 ist vorgeschrieben. Für σ können wir vorläufig einen Schätzwert benutzen, den wir aus schon akzeptierten Stichproben gewinnen können.

• *Nous savons par expérience et par conséquent nous l'acceptons aussi maintenant: $\bar{X} = \mu$ satisfait une distribution normale avec l'écart-type σ connu. Cette supposition est cependant seulement raisonnable dans un certain domaine parce que le domaine de définition de la répartition normale est $(-\infty, \infty) = \mathbb{R}$ tandis que le domaine, dans lequel on peut pratiquement trouver des valeurs de mesures x_i et par conséquent les moyennes $\bar{x}$, est $(0, M_{\max})$. $M_{\max}$ dépend des possibilités pratiques. Des valeurs de mesure négatives n'existent pas. μ_0 est prescrit, pour σ nous*

pouvons utiliser provisoirement une valeur estimative que nous pouvons obtenir d'échantillons déjà acceptés.

- (b) Die folgende Forderung oder Hypothese H_0 sei erfüllt: Der Sollwert μ_0 wird während der ganzen Produktion eingehalten, d.h. $H_0 : \mu = \mu_0$. Die Alternative dazu wäre $H_1 : \mu \neq \mu_0$.
- *L'hypothèse suivante H_0 soit satisfaite: La valeur prescrite μ_0 est respectée strictement pendant toute la production, ç.v.d. $H_0 : \mu = \mu_0$. L'alternative à cette hypothèse serait $H_1 : \mu \neq \mu_0$.*

Definition: • **Définition:**

H_0 heisst **Nullhypothese**, H_1 heisst **Alternativhypothese**.

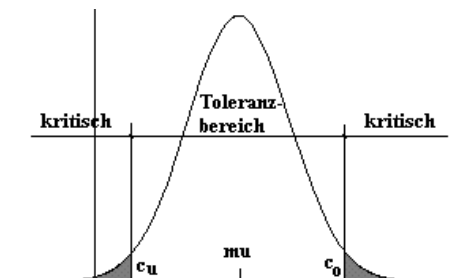
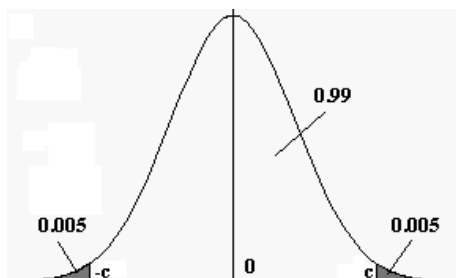
• H_0 s'appelle **hypothèse nulle**, H_1 s'appelle **hypothèse alternative**.

2. In regelmässigen Zeitabständen sollen der Produktion Stichproben von gleichbleibendem Umfang n entnommen werden. Daraus wird der Stichprobenmittelwert $\bar{x}$ bestimmt $\leadsto$ **Messexperiment**.
 - *Selon des intervalles de temps réguliers, on prélève à la production des échantillons de la taille invariable n . Après on en calcule la moyenne d'échantillon $\bar{x} \leadsto$ **expérience de mesure**.*
3. Weiter hat uns die Erfahrung gelehrt, dass es vernünftig ist die Produktion zu korrigieren, wenn die Wahrscheinlichkeit, dass die Alternativhypothese $H_1 : \mu \neq \mu_0$ eintritt, $\alpha = 0.01$ oder 1% übersteigt. (Wenn α erreicht ist, so bedeutet das etwas (fz. signifie!)! Man hat das Signifikanzniveau erreicht.) Dies ist so in Ordnung, denn bisher hat die Produktion sich etwa nach dem Gesetz der Normalverteilung verhalten, wobei 1% Ausschuss akzeptiert worden ist. Wir betrachten dazu die Standardnormalverteilung $N(0, 1)$ (links) im Vergleich zu $N(\mu, \sigma^2)$ (rechts):
 - *En outre nous avons appris de l'expérience qu'il est raisonnable de corriger la production, si la probabilité que l'hypothèse alternative $H_1 : \mu \neq \mu_0$ soit réalisée dépasse $\alpha = 0.01$ ou 1%. (Si α est atteint, ça signifie donc quelque chose! On a atteint le niveau de signification.) Ainsi c'est en ordre, car jusqu'à présent la production a montré environ la loi de la répartition normale et on a accepté 1% de rebut. Pour comprendre mieux nous considérons la distribution normale $N(0, 1)$ (à gauche) comparée avec $N(\mu, \sigma^2)$ (à droite):*

Definition: • **Définition:**

α heisst **Signifikanzniveau**. Die Zahl bezeichnet eine Grenze, an der meist nach Erfahrung etwas Bedeutendes geschieht.

• α s'appelle **niveau de signification**. Le nombre marque une limite à laquelle il se passe quelque chose d'important, dans la plupart des cas d'après l'expérience.



Der Übergang von der normalverteilten Zufallsvariablen $\bar{X}$ zur Standardisierten U ist durch die folgende Transformation gegeben:

• *On passe de la variable aléatoire distribuée de façon normale $\bar{X}$ à la variable aléatoire standardisée*

U par la transformation suivante:

$$U = \sqrt{n} \cdot \frac{\bar{X} - \mu_0}{\sigma}$$

In der standardisierten Form lautet die Bedingung mit α :

- Utilisant α , la condition est la suivante dans la forme standardisée:

$$P(-c \leq U \leq c) \leq 1 - \alpha = 1 - 0.01 = 0.99, \quad \pm c = \sqrt{n} \cdot \frac{c_{o,u}^* - \mu_0}{\sigma}, \quad c_{o,u}^* = \mu + \frac{\pm c \cdot \sigma}{\sqrt{n}}$$

4. Es gilt: • il vaut: $P(-c \leq U \leq c) = \frac{1}{\sqrt{2\pi}} \int_{-c}^c e^{-\frac{x^2}{2}} dx \leq 0.99 \leadsto c = ?$

Man kann auch rechnen: • On peut aussi déduire:

$$P(-c \leq U \leq c) = \Phi(c) - \Phi(-c) = \Phi(c) - (1 - \Phi(c)) = 2\Phi(c) - 1 = 2 \frac{1}{\sqrt{2\pi}} \int_{-\infty}^c e^{-\frac{x^2}{2}} dx$$

Mit Mathematica (oder Tabelle) findet man:

- on trouve à l'aide de Mathematica (ou dans un tableau):

`FindRoot[1/(Sqrt[2 Pi]) Integrate[E^(-x^2/2), {x, -c, c}] == 0.99, {c, 1}]`

$$\leadsto 2.57582 \leq c \leq 2.57583, \quad c \approx 2.5758 \Rightarrow c_o^* \approx \mu + \frac{2.5758 \cdot \sigma}{\sqrt{n}},$$

$$[c_u^*, c_o^*] \approx \left[\mu - \frac{2.5758 \cdot \sigma}{\sqrt{n}}, \mu + \frac{2.5758 \cdot \sigma}{\sqrt{n}} \right]$$

5. Sobald $\bar{x}$ aus dem Bereich $[c_u^*, c_o^*]$ herausläuft, muss also eingeschritten werden! Die Abweichung von $\bar{x}$ wird dann nicht mehr als zufallsbedingt eingestuft und toleriert, sondern als signifikant und nicht tolerierbar.

- Dès que $\bar{x}$ sort de l'intervalle $[c_u^*, c_o^*]$, il faut réagir! L'écart de $\bar{x}$ n'est plus classé comme aléatoire et donc n'est plus toléré. Elle est regardée comme significative et intolérable.

6. Oft wählt man zu den **kritischen Grenzen** $c_{\alpha,u,o}^* = c_{0.01,u,o}^*$ auch **Warn Grenzen** $d_{\beta,u,o}^* = d_{0.05,u,o}^*$. Dabei macht man für d zu $\beta = 0.05$ die entsprechende Rechnung wie für c zu $\alpha = 0.01$.

- Souvent on choisit pour les **limites critiques** $c_{\alpha,u,o}^* = c_{0.01,u,o}^*$ des **limites d'avertissement** $d_{\beta,u,o}^* = d_{0.05,u,o}^*$. Le calcul pour d avec $\beta = 0.05$ correspond au calcul pour c avec $\alpha = 0.01$.

$$\leadsto [d_u^*, d_o^*] \approx \left[\mu - \frac{1.95996 \cdot \sigma}{\sqrt{n}}, \mu + \frac{1.95996 \cdot \sigma}{\sqrt{n}} \right]$$

7. Zahlenbeispiel: • Exemple avec des nombres:

Sei • Soit $\mu = 120$, $\sigma = 0.01$, $n = \dots$

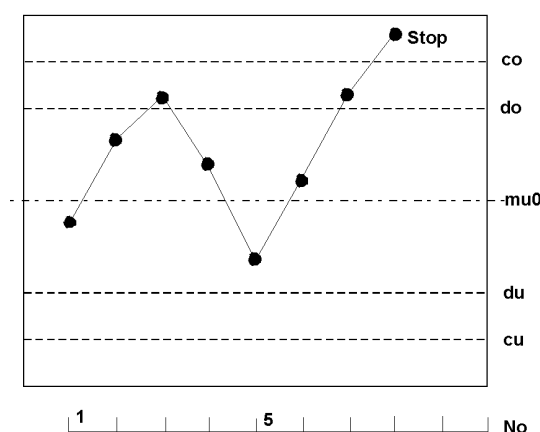
$c_{0.01,u,o}$	$n = 1$	$n = 2$	$n = 10$	$n = 20$	$n = 100$
c_u	119.974	119.982	119.992	119.994	119.997
c_o	120.026	120.018	120.008	120.006	120.003
$d_{0.05,u,o}$	$n = 1$	$n = 2$	$n = 10$	$n = 20$	$n = 100$
d_u	119.980	119.986	119.994	119.996	119.998
d_o	120.020	120.014	120.006	120.004	120.002

5.2.3 Überwachungstechnik mittels Kontrollkarten — Technique de surveillance à l'aide de cartes de contrôle

Es ist oft üblich, die Resultate der Stichprobenerhebungen graphisch übersichtlich und anschaulich darzustellen. Dazu verwendet man in der Industrie **Kontrollkarten** diverser Art. Exemplarisch geben wir die Kontrollkarte zu einer Serie von Stichproben zu obigem Beispiel.

• *Souvent il est usuel de représenter les résultats des échantillonnages graphiquement et clairement. Dans l'industrie on utilise des différentes cartes de contrôle. Nous donnons exemplairement une carte de contrôle concernant une série d'échantillons concernant l'exemple susdit.*

$$\mu = 120, \sigma = 0.01, n = 10, \alpha = 0.01, \beta = 0.05, c_{\alpha,u,o}^*, d_{\beta,u,o}^*$$



5.3 SQC2: Annahmekontrolle — SQC2: Contrôle d'acceptation

Spez. franz.. • *Contrôle d'acceptation* $\leadsto$ *de réception etc..*

Stichproben und Randomisierung — Echantillons et randomisation

Problem: • Problème:

Als Kunden erhalten wir eine grössere Lieferung einer Ware. Man denke z. B. an Zahnräder. Da wir auf die Einhaltung der vertraglich abgesicherten Toleranzen angewiesen sind um die Teile problemlos einbauen und die Funktion garantieren zu können, müssen wir die Einhaltung der Toleranzen der Teile prüfen. Wir führen daher eine **Annahmekontrolle** durch, um auf Grund der Verträge zu entscheiden, ob die Lieferung angenommen werden kann oder nicht. Falls es sich z.B. um die Prüfung der chemischen Zusammensetzung handelt (z.B. Prüfung der Bleifreiheit der Glasuren von Keramik-Essgeschirr), ist oft zur Prüfung eine Zerstörung des Stückes notwendig. Weil aber auch sonst eine Prüfung aller Teile zu hohe Kosten verursacht, beschränken wir uns auf Stichproben und mathematische Methoden, um daraus Rückschlüsse auf die Grundgesamtheit zu ziehen. Stichproben lassen sich exakter kontrollieren als grosse Stückzahlen, wo die Routine auf Kosten der Sorgfalt geht. Überlegt man sich die Konsequenzen eines falschen Entscheides betreffend Annahme und Weiterverarbeitung einer Lieferung, so merkt man, dass davon das überleben einer Firma abhängen kann.

• *Etant clients nous recevons une grande livraison d'une marchandise. On pense par exemple à des roues dentées. Comme nous sommes obligés de respecter les tolérances assurées par contrat pour pouvoir monter les pièces sans problèmes et de pouvoir garantir le fonctionnement, nous devons examiner les tolérances des pièces. Nous effectuons par conséquent un **contrôle d'acceptation** pour pouvoir décider sur la base des contrats si la livraison peut être accepté ou non. S'il s'agit par exemple d'examiner de composition chimique (p.ex. examiner si les vernis (glaçure, couverte) des services de table contiennent du plomb), une destruction de la pièce lors de l'examen est souvent nécessaire. Comme d'autre part*

un examen de toutes les parties cause des frais très hauts, nous nous limitons aux échantillons et les méthodes mathématiques pour tirer des conclusions valables pour la totalité de l'ensemble fondamental. On peut contrôler les échantillons plus exactement qu'un grand nombre de pièces, où la routine remplace le soin. Si on considère les conséquences d'une fausse décision à propos de l'acceptation et de l'utilisation d'une livraison dans la production, on s'aperçoit la survie d'une entreprise en peut dépendre.

Die Modalitäten einer **Annahmekontrolle** des Kunden müssen entsprechend den Abmachungen in einem Prüfplan fixiert sein. Die genau gleiche Kontrolle kann auch der Lieferant für sich durchführen. Hier heisst sie dann **Ausgangskontrolle**, **Endkontrolle** u.s.w..

• *Les modalités d'un **contrôle d'acceptation** du client doivent être fixées selon les conventions dans le plan de contrôle. Le fournisseur peut aussi effectuer exactement le même contrôle pour lui-même. Ici, elle s'appelle le **contrôle de sortie**, **contrôle final** etc..*

Bei der Erhebung der Stichprobe spielt die Zufälligkeit der Stückauswahl eine wesentliche Rolle, um systematische Verfälschungen und Bevorzugungen auszuschliessen. Es ist daher sehr gefährlich, gedankenlos oder blind irgendwelche ersten Stücke auszuwählen. Die Auswahl muss nach einem **Randomisierungsverfahren** geschehen. Z.B. kann man den Stücken Nummern zuordnen und dann diejenigen Nummern wählen, die in einer Randomisierungsliste als Zufallszahlen stehen oder von einem Pseudo-Zufallsgenerator generiert worden sind. Zum Problem der Generierung von Zufallszahlen muss des engen Rahmens wegen auf die Literatur verwiesen werden. (Möglichkeit: Nehme Ziffern aus der Dezimalbruchentwicklung von π .)

• *Au tirage de l'échantillon, le hasard du choix des pièces joue un rôle considérable pour exclure des falsifications systématiques et des préférences. Par conséquent il est très dangereux de choisir aveuglement quelques premières pièces. Le choix doit être fait d'après un système de **randomiser** le choix. P.ex. on peut attribuer un numéro à chaque pièce. Alors on choisit les numéros qu'on trouve dans une liste randomisée ou qui ont été générés par un pseudo-générateur de nombres aléatoires. Quant au problème de générer des nombres aléatoires, le lecteur est renvoyé à la littérature à cause de notre cadre étroit. (Possibilité: Prendre des chiffres du nombre π développé en fraction décimale.)*

5.3.1 Beispiel hypergeometrische Verteilung, Urnenmodell — Exemple répartition hypergéométrique, modèle d'urne

Oben haben wir gesehen, dass man generell unterscheiden kann zwischen qualitativen Kontrollen (z.B. in Ordnung, nicht in Ordnung) und quantitativen Kontrollen (Messung von Variablen).

• *En haut on a vu qu'on peut généralement distinguer entre les contrôles qualitatifs ou contrôles par attributs (p.ex. en ordre, pas en ordre) et les contrôles quantitatifs ou contrôles par variables (mesurage de variables).*

Wir betrachten nun hier ein Beispiel aus der Klasse der qualitativen Annahmekontrollen. Gegeben sei eine Warenlieferung mit N Teilen. Davon seien M Teile defekt und daher $N - M$ Teile in Ordnung. Statt eine Warenlieferung zu betrachten, können wir auch an N Kugeln in einer Urne denken: M schwarze und $N - M$ weisse.

• *Maintenant nous considérons un exemple de la classe des contrôles d'acceptation qualitatifs. Soit donné une livraison de marchandises à N parties. M parties en sont défectueuses et par conséquent $N - M$ parties sont en ordre. Au lieu de considérer une livraison de marchandises, nous pouvons aussi penser à N boules dans une urne: M sont noires et $N - M$ sont blanches.*

Aus der Lieferung resp. aus der Urne nehmen wir eine Zufallsstichprobe von n Stück ohne zurücklegen, $n \leq N$. Von früher wissen wir, dass die möglichen Resultate dieses Experiments ein Modell für die hypergeometrische Verteilung bilden. Sei $X = m$ die Anzahl der gezogenen defekten Stücke — oder beim Urnenmodell die Anzahl gezogener schwarzer Kugeln. Sei $m \in \{1, 2, \dots, n\}$. Dann wissen wir von der

hypergeometrische Verteilung:

• De la livraison resp. de l'urne, nous tirons un échantillon aléatoire consistant en n pièces et sans les remettre, $n \leq N$. De la théorie nous savons que les résultats possibles de cette expérience forment un modèle pour la répartition hypergéométrique. Soit $X = m$ le nombre des pièces défectueuses qu'on vient de tirer — ou lors du modèle d'urne le nombre de boules noires tirées. Sois $m \in \{1, 2, \dots, n\}$. Alors nous connaissons de la répartition hypergéométrique:

$$p_d = \frac{M}{N}, \quad M = p_d \cdot N, \quad p_m = P(X = m) = \frac{\binom{M}{m} \cdot \binom{N-M}{n-m}}{\binom{N}{n}}$$

Bsp.: • Exemple: $N = 120, M = 4, n = 3$

$$\leadsto p_0 = P(X = 0) = \frac{\binom{4}{0} \cdot \binom{120-4}{3-0}}{\binom{120}{3}} \approx 0.902507,$$

$$p_1 = P(X = 1) \approx 0.0950007, \quad p_2 = P(X = 2) \approx 0.00247828, \quad p_3 = P(X = 3) \approx 0.000014243$$

$$\Rightarrow P(X \leq 3) = P(X = 0) + P(X = 1) + P(X = 2) + P(X = 3) = 1$$

5.3.2 Annahmevertrag und Prüfplan — Contrat d'acceptation et plan d'échantillonnage

Produzent und Konsument resp. Lieferant und Kunde müssen sich nun über die Annahmebedingung der Lieferung einig sein. Die Annahmebedingung muss in einem Vertrag (Lieferungsvertrag, Annahmevertrag) geregelt sein. Die Beschreibung der Annahmebedingungen mündet in einen Prüfungsplan (**Stichprobenprüfplan**), in dem die Details der Annahmeprüfung definiert sind.

• Le producteur et le consommateur resp. le fournisseur et le client doivent être d'accord en ce qui concerne les conditions d'acceptation de la livraison. Les conditions d'acceptation doivent être réglées dans un contrat, le contrat de livraison ou le contrat d'acceptation. La description des conditions d'acceptation débouchent dans un plan d'examen, le **plan d'échantillonnage** dans lequel les détails de l'examen d'acceptation sont définis.

Als Beispiel, das zur oben beschriebenen Situation passt, studieren wir einen einfachen Stichprobenprüfplan (**SSP, simple sampling plan**). Dann definieren wir:

• Nous étudions comme exemple, qui va avec la situation décrite en haut, un plan d'échantillonnage simple (**SSP, simple sampling plan**).

Vor.: • Hyp.:

Sei X die Anzahl defekter Teile resp. schwarzer Kugeln in der gezogenen Stichprobe ($0 \leq X = x \leq n$).

• Soit X le nombre de pièces défectueuses resp. de boules noires dans l'échantillon tiré ($0 \leq X = x \leq n$).

Dann definieren wir:

• Alors nous définissons:

Definition: • Définition:

Ein **einfacher Stichprobenprüfplan** ist gegeben durch:

- **Un plan d'échantillonnage simple** est donné par:
 1. Stichprobenumfang n
 - *La taille de l'échantillon* n
 2. Eine Annahmezahl (**kritische Zahl**) $c \in \mathbb{R}_0^+$
 - *Un nombre d'acceptation (nombre critique)* $c \in \mathbb{R}_0^+$
 3. Eine Entscheidungsregel:
 - *Une règle de décision:*

$X \leq c \rightsquigarrow$	Annahme • <i>acceptation</i>
$X > c \rightsquigarrow$	Ablehnung • <i>lot refusé</i>

Problem: • Problème:

Allgemein kennt man bei einer Lieferung $p_d = \frac{M}{N} \in [0, 1]$ nicht. p_d ist daher ein unbekannter Wert einer Zufallsvariablen Y . Es ist also notwendig, die Annahmewahrscheinlichkeit der Lieferung in Funktion von p_d in Modellen zu untersuchen, um die theoretisch gewonnenen Erkenntnisse praktisch dann kostensenkend verwenden zu können.

• *Généralement on ne connaît pas $p_d = \frac{M}{N} \in [0, 1]$ lors d'une livraison. Par conséquent p_d est une valeur inconnue d'une variable aléatoire Y . Ainsi il est nécessaire d'examiner la probabilité d'acceptation de la livraison en fonction de p_d par des modèles, pour pouvoir appliquer pratiquement les connaissances acquises de façon théorique à l'intention de diminuer les frais.*

Im oben besprochenen Beispiel (hypergeometrische Verteilung, Urnenmodell) ist $X = m$ und $M = p_d \cdot N$. Daher gilt für $p_m = P(X = m)$:

• *Dans l'exemple discuté en haut (répartition hypergéométrique, modèle d'urne) il était $X = m$ et $M = p_d \cdot N$. Il vaut donc pour $p_m = P(X = m)$:*

$$p_m = P(X = m) = \frac{\binom{p_d \cdot N}{m} \cdot \binom{N - p_d \cdot N}{n - m}}{\binom{N}{n}} = \frac{\binom{p_d \cdot N}{m} \cdot \binom{N(1 - p_d)}{n - m}}{\binom{N}{n}}$$

$$\Rightarrow P(X = m \leq c) = \sum_{i, m_i \leq c} P(X = m_i) = \sum_{m=0}^{m \leq c} P(X = m) = \sum_{m=0}^{m \leq c} \frac{\binom{p_d \cdot N}{m} \cdot \binom{N(1 - p_d)}{n - m}}{\binom{N}{n}}$$

$P(X = m \leq c)$ ist die Wahrscheinlichkeit, maximal c defekte Stücke resp. schwarze Kugeln in der Stichprobe zu haben. Dabei ist c fix, $p_d = \frac{M}{N}$ jedoch eine unbekannte Realisation von Y . $P(X = m \leq c)$ ist somit eine Funktion von Y mit den bekannten Parametern N, n, c . $P(X = m \leq c)$ ist die **Annahmewahrscheinlichkeit** $L_{N,n,c}$ in Abhängigkeit von Y :

• *$P(X = m \leq c)$ est la probabilité de trouver au maximum c pièces défectueuses resp. c boules noires dans l'échantillon. c est fixé par contrat, mais par contre $p_d = \frac{M}{N}$ est une réalisation inconnue de Y . Donc $P(X = m \leq c)$ est une fonction de Y aux paramètres connus N, n, c . $P(X = m \leq c)$ est la **probabilité d'acceptation** $L_{N,n,c}$ qui dépend de Y :*

$$L_{N,n,c}(Y) := P(X = m \leq c) = \sum_{m=0}^{m \leq c} \frac{\binom{Y \cdot N}{m} \cdot \binom{N(1 - Y)}{n - m}}{\binom{N}{n}}, \quad Y \in [0, 1]$$

Generell gehört zu einem einfachen Versuchsplan eine Funktion $L_c(Y) := P(X \leq c)$, $Y = p_d \in [0, 1]$.

• *Généralement on a toujours une fonction $L_c(Y) := P(X \leq c)$, $Y = p_d \in [0, 1]$ liée à un plan d'échantillonnage.*

Definition: • **Définition:**

$L_c(Y) := P(X \leq c)$, $Y = p_d \in [0, 1]$

heisst **Annahmewahrscheinlichkeitsfunktion** (der Lieferung oder des Loses) durch den Konsumenten

• s'appelle **fonction de probabilité d'acceptation** (du lot par le consommateur)

Der Graph von L_c heisst **Effizienzkurve** oder **Operationscharakteristik** des einfachen Versuchsplanes

• Le graphe de L_c s'appelle **courbe d'efficacité** du plan d'échantillonnage simple

Bemerkung: • **Remarque:**

$\leadsto$ Engl. • *Angl. operating characteristic, OC-curve*

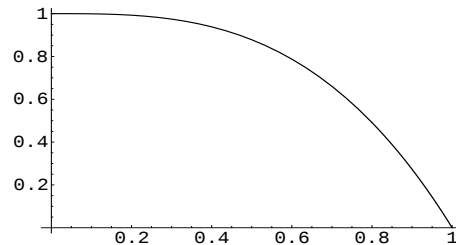
In obigem Beispiel sei z.B. $c = 2$: • *Dans l'exemple en haut c soit p.ex. $c = 2$:*

$$L_{120,3,c=2}(Y) := P(X = m \leq 2) = \sum_{m=0}^{m \leq 2} \frac{\binom{Y \cdot 120}{m} \cdot \binom{120(1-Y)}{3-m}}{\binom{120}{3}}$$

$$= \frac{\binom{Y \cdot 120}{0} \cdot \binom{120(1-Y)}{3}}{\binom{120}{3}} + \frac{\binom{Y \cdot 120}{1} \cdot \binom{120(1-Y)}{2}}{\binom{120}{3}} + \frac{\binom{Y \cdot 120}{2} \cdot \binom{120(1-Y)}{1}}{\binom{120}{3}}$$

Speziell: • **Spécialement:**

$$L_{120,3,c=2}(0.9) \approx 0.273052$$



Eigenschaften von L_c : • **Qualités de L_c :**

$L_c(0) = P(X \leq c)_{|M=0}$ bedeutet, dass kein defektes Stück da ist und daher die Annahme sicher ist $\leadsto Y = p_d = 1$

Wenn die Anzahl der defekten Stücke steigt, so wird auch die Chance, dass bei einer Stichprobe defekte Stücke gefunden werden, grösser. Die Annahmewahrscheinlichkeit wird damit also kleiner. Somit muss $L_c(Y)$ mit Y monoton fallend sein.

$L_c(1) = P(X \leq c)_{|M=N}$ bedeutet, dass nur defekte Stücke da sind und daher die Annahme sicher verweigert wird $\leadsto Y = p_d = 0$

• $L_c(0) = P(X \leq c)_{|M=0}$ signifie qu'il n'y a pas de pièces défectueuses et par conséquent que l'acceptation est sûre $\leadsto Y = p_d = 1$

Si le nombre des pièces défectueuses augmente, la chance de trouver des pièces défectueuses lors d'un échantillonnage augmente aussi. La probabilité d'acceptation diminue donc. Par conséquent $L_c(Y)$ doit décroître de façon monotone avec Y .

$L_c(1) = P(X \leq c)_{|M=N}$ signifie qu'il n'y a que des pièces défectueuses et que l'acceptation sera donc

refusée $\rightsquigarrow Y = p_d = 0$

Eigenschaften: $L_c(0) = 1, L_c(1) = 0, D_L = [0, 1,]$
• Qualités: $L_c(Y)$ monoton fallend • *décroissant de façon monotone*

5.3.3 Das Beispiel Binomialmodell, Urnenmodell — L'exemple modèle binomial, modèle d'urne

Oben hatten wir eine Urne und die Parameter M, N, n, m . Das Experiment ist ohne zurücklegen ausgeführt worden. Grundlage war eine hypergeometrische Verteilung. Nun wollen wir das Experiment mit zurücklegen ausführen. Das führt uns auf eine Binomialverteilung

• *En haut nous avons eu une urne et les paramètres M, N, n, m . Nous avons fait l'expérience sans remettre les pièces. La base était donc une répartition hypergéométrique. Maintenant nous voulons exécuter l'expérience en remettant la pièce chaque fois. Cela nous mène à une distribution binomiale.*

Sei • *Soit* $H = \{1, 2, \dots, M, M+1, \dots, N\}$ Wir können davon ausgehen, dass wir ein geordnetes n -Tupel aus der geordneten Menge H auswählen. Dann finden wir:

• *Nous pouvons considérer le fait que nous choisissons un système ordonné à n éléments de l'ensemble H . Alors nous trouvons:*

Totale Anzahl Auswahlmöglichkeiten = N^n • *Nombre total de possibilités de choisir = N^n*

Totale Anzahl Möglichkeiten nur defekte Stücke zu wählen = M^m • *Nombre total de possibilités de choisir que des pièces défectueuses = M^m*

Totale Anzahl Möglichkeiten nur gute Stücke zu wählen = $(N - M)^{n-m}$ • *Nombre total de possibilités de choisir que des pièces bonnes = $(N - M)^{n-m}$*

Anzahl Möglichkeiten in der Stichprobe Plätze an die defekten Stücke zu verteilen: = $\binom{n}{m}$ • *Nombre de possibilités dans l'échantillonnage de distribuer des places aux pièces défectueuses = $\binom{n}{m}$*

Die Anzahl Möglichkeiten, m defekte und zugleich $n - m$ gute Stücke zu wählen und dann noch Plätze auf die defekten zu verteilen, d.h. die defekten und die guten zu vermischen, ist gleich der Anzahl Stichproben mit m defekten Stücken bei zurücklegen. Dividiert man durch die Gesamtzahl N^n der Möglichkeiten, so erhält man die Laplace-Wahrscheinlichkeit $P(X = m)$: • *Le nombre de possibilités de choisir m pièces défectueuses et au même temps $n - m$ pièces bonnes et en plus de distribuer les places aux pièces défectueuses, c.v.d. de mélanger les pièces défectueuses et les autres, ceci nous donne le nombre d'échantillonnages à m pièces défectueuses en remettant la pièce chaque fois après l'avoir tirée. Si on divise encore par le nombre total de possibilités de choisir N^n , on obtient la probabilité de Laplace $P(X = m)$:*

$$\frac{M}{N} := p \Rightarrow P(X = m) = \binom{n}{m} \cdot \frac{M^m \cdot (N - M)^{n-m}}{N^n} = \binom{n}{m} \cdot \frac{M^m}{N^m} \cdot \frac{(N - M)^{n-m}}{N^{n-m}} = \binom{n}{m} \cdot p^m \cdot (1-p)^{n-m}$$

Bekanntlich haben wir hier eine Binomialverteilung.

• *Comme nous savons, nous avons ici une répartition binomiale.*

Folgender Zusammenhang rechtfertigt den Namen: • *La loi suivante justifie le nom:*

$$\sum_{m=1}^n p_m = \sum_{m=1}^n \binom{n}{m} \cdot p^m \cdot (1-p)^{n-m} = (1 + (1-p))^n = 1^n = 1$$

Vor.: • **Hyp.:** Sei • Soit $n \ll N$, $m \in \{0, 1, \dots, n\}$, $p = \frac{M}{N}$, $N \rightarrow \infty$

$\leadsto$ Dann kann man die hypergeometrische Verteilung durch die Binomialverteilung approximieren (praktisch z.B. für $10n \leq N$):

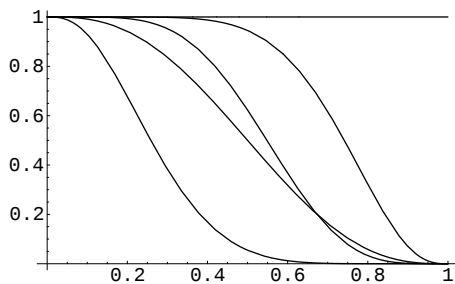
• Alors on peut approximer la répartition hypergéométrique par la répartition binomiale (en pratique pour $10n \leq N$):

$$\lim_{N \rightarrow \infty} \frac{\binom{M}{m} \cdot \binom{N-M}{n-m}}{\binom{N}{n}} = \binom{n}{m} \cdot p^m \cdot (1-p)^{n-m}$$

In diesem Fall gilt für die Annahmewahrscheinlichkeitsfunktion:

• Dans ce cas il vaut pour la fonction de probabilité d'acceptation:

$$L_{n,c}(Y) = P(X \leq c) = \sum_{m=0}^{m \leq c} \binom{n}{m} \cdot Y^m \cdot (1-Y)^{n-m}$$



Links Kurven für: • A gauche courbes pour: $(n, c) = (5, 2), (10, 2), (10, 5), (10, 7), (5, 5)$

Beobachtung: Je grösser c , desto weiter rückt die Kurve nach rechts. Je grösser n , desto weiter rückt die Kurve nach unten.

• Observation: Un c plus grand pousse la courbe à droite. Un n plus grand pousse la courbe vers le bas.

5.3.4 Produzentenrisiko, Konsumentenrisiko — Risque du producteur, risque du consommateur

$L_{n,c}(Y) = L_{n,c}(p_d) = P(X \leq c)_{|Y=p_d}$ gibt bei den Parametern n und c die Wahrscheinlichkeit $P(X \leq c)$ in Abhängigkeit von $Y = p_d$, d.h. bei unbekanntem p_d . Da p_d unbekannt ist und trotzdem ein Entscheid gefällt werden muss ($X \leq c \leadsto$ Annahme, $X > c \leadsto$ Ablehnung), so kann es vorkommen, dass der Entscheid bei falsch angenommenem p_d zu ungunsten des Produzenten oder auch zu ungunsten des Konsumenten ausfällt.

• Utilisant les paramètres n et c , $L_{n,c}(Y) = L_{n,c}(p_d) = P(X \leq c)_{|Y=p_d}$ donne la probabilité $P(X \leq c)$ dépendante de $Y = p_d$ si p_d est inconnu. Comme p_d est inconnu et lorsqu'on doit décider quand-même ($X \leq c \leadsto$ acceptation, $X > c \leadsto$ lot refusé), il peut arriver que la décision est au désavantage du producteur ou du consommateur, parce que p_d à été mal estimé.

Wird z.B. die **Ablehnungsgrenze** oder Schätzung $Y = p_\beta$ des effektiven Ausschussanteils p_d vom Konsumenten zu tief angesetzt ($Y = p_\beta < p_d$), so kann es vorkommen, dass die **Lieferung nach Prüfplan akzeptiert** wird, obwohl sie einen nicht **tolerierbaren Ausschussanteil** aufweist. Denn die Annahmewahrscheinlichkeit könnte dann zu gross werden.

• Si la limite de refus ou l'estimation $Y = p_\beta$ de la portion de rebut effectif p_d est fixée trop bas par le consommateur ($Y = p_\beta < p_d$), il peut arriver, que la livraison est acceptée d'après le plan d'acceptation, bien qu'elle contienne une portion de rebut intolérable. Car la probabilité d'acceptation pourrait devenir trop grande.

Definition: • **Définition:**

Die Annahme eines Loses durch den Konsumenten trotz zuviel Ausschuss heisst **Fehler 2. Art**. Die dazugehörige Annahmewahrscheinlichkeit $\beta = L_{n,c}(p_\beta) = P(X \leq c)$ heisst **Konsumentenrisiko**. p_β heisst **Ablehnungsgrenze**.

• *L'acceptation d'un lot par le consommateur malgré trop de rebut s'appelle **erreur de la 2-ième sorte**. La probabilité d'acceptation nécessaire $\beta = L_{n,c}(p_\beta) = P(X \leq c)$ s'appelle le **risque de consommation**. p_β s'appelle **limite de refus**.*

(Hit! Konsumentenrisiko wird von der Maschine mit "Le risque de canard de consommation" übersetzt.

• *Hit! A réfléchir ...)*

Entsprechend definieren wir: • *Au fur et à mesure nous définissons:*

Definition: • **Définition:**

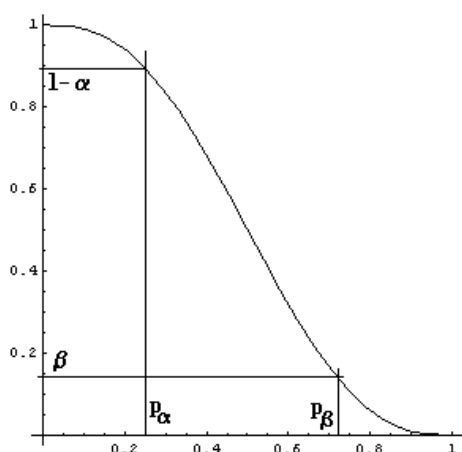
Die Ablehnung eines Loses obwohl es in Ordnung ist und der Ausschuss tolerierbar ist, heisst Fehler 1. Art. Die zugehörige Wahrscheinlichkeit $\alpha = P(X > c) = 1 - P(X \leq c) = 1 - L_{n,c}(p_\alpha)$ heisst **Produzentenrisiko**. p_α heisst **Annahmegrenze**.

• *Le refus d'un lot bien qu'il soit en ordre et la portion de rebut soit tolérable s'appelle **erreur de la 1-ère sorte**. La probabilité liée $\alpha = P(X > c) = 1 - P(X \leq c) = 1 - L_{n,c}(p_\alpha)$ s'appelle **risque du producteur**. p_α s'appelle **limite d'acceptation**.*

Da α und β Fehlerwahrscheinlichkeiten sind, muss der Prüfplan so eingerichtet werden (Vertrag!), dass die beiden Grössen möglichst klein werden. Aus dem Diagramm ist ersichtlich, dass das erreicht werden kann, wenn die Kurve möglichst steil gewählt wird (Modellierung!) oder wenn p_α möglichst rechts und p_β möglichst links liegt.

• *Comme α et β sont des probabilités d'erreurs, il faut construire le plan d'échantillonnage (contrat!) de façon que les deux grandeurs deviennent aussi petites que possible. Comme on voit dans le diagramme, il est évident qu'on atteint ce fait si on choisit une courbe très raide (modélisation!) et si p_α est situé le plus que possible à gauche et p_β à droite.*

Man bedenke, dass grosse n mehr kosten! • *Il faut y penser que les n plus grands coûtent davantage.*

**Idee:** • **Idée:**

Sei emp = empirisch

• *Soit emp = empirique*

Wähle • *Choisir*

$$\left(\frac{M}{N}\right)_{emp} \leq p_\alpha \leq \left(\frac{M}{N}\right)_{emp} + \varepsilon$$

Wird hier M während der Produktion grösser, so sollte man das merken!

• *Si M devient plus grand pendant la production, on devrait maintenant le constater!*

Bemerkung: • **Remarque:**

Stichprobenprüfung: $\leadsto p_\alpha$ heisst auch **annehmbare Qualitätslage**, **Gutlage** oder **AQL** (acceptable quality level). p_β heisst auch **rückweisende Qualitätsgrenzlage**, **Schlechtlage** oder **LQ, LQL, RQL** (limiting quality level, rejectable quality level).

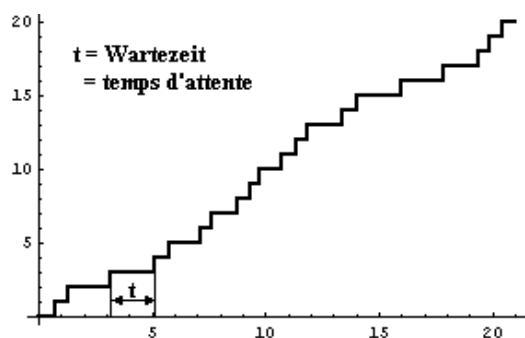
• *Echantillonnage*: $\leadsto p_\alpha$ s'appelle aussi la **limite de qualité acceptable** ou **AQL** (acceptable quality level). p_β s'appelle aussi la **limite de qualité tolérée** ou **LQ, LQL, RQL** (limiting quality level, rejectable quality level).

Normierungen: • *Normalisations:*

DIN ISO 2859, US Military Standard 105 D, DIN 40080, Philips-Standard, DIN ISO 3951, ISO 8423 etc.
 $1\% \leq \alpha\% \leq 12\%$, norm. $\beta\% \leq 10\%$, except. $\beta\% \leq 13\% \dots$ (Lit. Storm, Bibl. A12.)

5.4 Poisson-Prozess, Warteschlange — Files d'attente, processus de Poisson

5.4.1 Poisson-Prozess — Processus de Poisson



Sei $X(t)$ = Anzahl Zufallsereignisse von der Zeit 0 bis t

$\leadsto$ Die Zufallsvariable (stochastische Variable) ist zeitabhängig ($\leadsto$ Prozess) sowie diskret ($X(t) \in \mathbb{N}$).

• *Soit* $X(t)$ = nombre d'événements aléatoires du temps 0 jusqu'à t

$\leadsto$ La variable aléatoire (stochastique) dépend du temps ($\leadsto$ processus) et en plus elle est discrète ($X(t) \in \mathbb{N}$).

Definition: • **Définition:**

Erfüllt $X(t)$ obige Bedingungen, so reden wir von einem **stochastischen Prozess**.

• *Si* $X(t)$ satisfait les conditions susdites, nous parlons d'un **processus stochastique**.

Da $X(t)$ von Natur aus nicht abnehmen kann, gilt:

• *Comme par sa nature* $X(t)$ ne peut pas diminuer, il vaut:

Lemma: • **Lemme:** $X(t)$ ist monoton wachsend. • $X(t)$ s'accroît de façon monotone.

Vor.: • **Hyp.:**

Zur Vereinfachung gehen wir hier von der Situation mit der **Startbedingung** $X(0) = 0$ aus.

• *Pour simplifier le problème, nous partons ici de la situation de la condition initiale* $X(0) = 0$.

Sei • *Soit* $\Delta X(t_0, t) := X(t_0 + t) - X(t_0)$

$\leadsto$ Der Zuwachs $\Delta X(t_0, t)$ von X im Intervall $[t_0, t_0 + t]$ hängt von t_0 und t ab.

- L'augmentation $\Delta X(t_0, t)$ de X dans l'intervalle $[t_0, t_0 + t]$ dépend de t_0 et t .

Wir betrachten hier einen einfachen Fall: • Ici nous considérons un cas simple:

Definition: • **Définition:** Der Prozess $X(t)$ heisst **homogen**, falls $\Delta X(t_0, t)$ unabhängig ist von t_0 .
 • Le processus s'appelle **homogène**, si $\Delta X(t_0, t)$ est indépendant de t_0 .

Dann gilt: • Alors il vaut:

$$\forall t_i: \Delta X(t_0, t) = \Delta X(t_1, t) \Rightarrow \Delta X(t_0, t) = \Delta X(0, t) = X(0 + t) - X(0) = X(t) - 0 = X(t)$$

Lemma: • **Lemme:** **Vor.:** • **Hyp.:**

$X(t)$ homogen • homogène
 $\Delta X(t_0, t) := X(t_0 + t) - X(t_0)$

Beh.: • **Thè.:**

$$\Delta X(t_0, t) = \Delta X(0, t) = X(t)$$

Sei • Soit $p_m(t) := P(X(t) = m)$

- $\leadsto p_m(t)$ ist die Wahrscheinlichkeit, dass in $[0, t]$ genau m Ereignisse liegen.
 • $\leadsto p_m(t)$ est la probabilité que dans $[0, t]$ il y a exactement m événements.

Definition: • **Définition:** Zwei Ereignisse heissen **statistisch unabhängig**, wenn das Auftreten irgend eines der beiden Ereignisse keinen Einfluss auf das Auftreten des andern Ereignis hat.
 • Deux événements s'appellent **statistiquement indépendants**, si l'apparition de n'importe quel des deux événements n'a aucune influence sur l'apparition de l'autre des deux événements.

Wir betrachten folgende Situation: • Nous considérons la situation suivante:

$$t_0 < t_1 < t_2 < \dots < t_n, \Delta X_r := X(t_r) - X(t_{r-1}), r \in \{1, 2, 3, \dots, n\}$$

Die Intervalle $I_r = [t_r, t_{r+1}]$ haben nichts gemeinsam ausser den Grenzen, wenn es sich um Nachbarintervalle handelt. Daher wollen wir hier voraussetzen, dass die ΔX_r statistisch unabhängig sind.

- Les intervalles $I_r = [t_r, t_{r+1}]$ n'ont rien en commun sauf les bords s'il s'agit d'intervalles voisins. C'est pourquoi nous voulons présupposer ici que les ΔX_r soient statistiquement indépendants.

$S = \sum_{m=0}^{\infty} p_m(\Delta t)$ = Wahrscheinlichkeit, dass in Δt 0 oder 1 oder 2 oder ... oder unendlich viele Ereignisse auftreten. Jede Zahl von Ereignissen ist daher möglich $\leadsto S = 1$. $\leadsto$

- $S = \sum_{m=0}^{\infty} p_m(\Delta t)$ = probabilité qu'il y a dans Δt 0 ou 1 ou 2 ou ... ou un nombre infini d'événements. Chaque nombre d'événements est donc possible $\leadsto S = 1$. $\leadsto$

Lemma: • **Lemme:** $1 - p_0(\Delta t) - p_1(\Delta t) = \sum_{m=2}^{\infty} p_m(\Delta t) := S_2$

S_2 ist die Wahrscheinlichkeit, dass in Δt mehr als 2 Ereignisse auftreten. In der Praxis ist es oft vernünftig anzunehmen, dass für kleine Δt $S_2 \approx 0$ gilt. Das wollen wir hier auch tun:

• S_2 est la probabilité que dans Δt il y a plus de 2 événements qui se sont réalisés. En pratique, il est souvent raisonnable de présupposer que pour des Δt petits il vaut $S_2 \approx 0$. C'est ce que nous voulons faire aussi ici:

Vor.: • **Hyp.:**

$$S_2 = 1 - p_0(\Delta t) - p_1(\Delta t) = \sum_{m=2}^{\infty} p_m(\Delta t) \approx 0$$

Konsequenz: • **Conséquence:** $\leadsto p_0(\Delta t) \approx 1 - p_1(\Delta t)$

Weiter ist es in der Praxis oft vernünftig anzunehmen, dass die Wahrscheinlichkeit für ein Ereignis $p_1(\Delta t)$ und für kleine Δt proportional zur Zeit ist. Auch das wollen wir hier tun:

• En pratique il est en outre raisonnable de présupposer que la probabilité pour un événement $p_1(\Delta t)$ et pour des temps Δt petits soit proportionnel au temps. C'est ce que nous voulons faire ici:

Vor.: • **Hyp.:**

$$p_1(\Delta t) \approx a \cdot \Delta t, \quad \Delta t < \text{const.}, \quad a > 0$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

1. $X(0) = 0$
2. $X(t)$ homogen • *homogène*
3. ΔX_r ($r = 1, 2, \dots$) statist. unabh.: $\in A$ • *statist. indép.: $\in A$*
4. $0 < \Delta t \ll \text{const.} \Rightarrow p_0(\Delta t) \approx 1 - p_1(\Delta t)$
5. $0 < \Delta t \ll \text{const.} \Rightarrow p_1(\Delta t) \approx a \cdot \Delta t, \quad a > 0$
6. $S_2 = \sum_{m=2}^{\infty} p_m(\Delta t) = \varphi(\Delta t) \cdot (\Delta t)^2, \quad |\varphi(\Delta t)| < \text{const.}$

Beh.: • **Thè.:**

$X(t)$ ist poissonverteilt mit $\lambda = a \cdot t$:

• $X(t)$ est réparti d'après Poisson avec $\lambda = a \cdot t$:

$$p_m(t) = P(X(t) = m) = \frac{(a t)^m}{m!} e^{-a t} = \frac{\lambda^m}{m!} e^{-\lambda}$$

Bemerkung: • **Remarque:**

Modell für: $X(t)$ = Anzahl Kunden, die an einem Bedienungsschalter eintreffen oder $X(t)$ = Anzahl Ersetzungen von Bauteilen mit identischer Lebenserwartung und gleicher Betriebszeit.

• *Modèle pour: $X(t)$ = nombre de clients qui arrivent à un bouton de service ou $X(t)$ = nombre de remplacements de composants avec l'espérance de vie identique et la même durée de fonctionnement.*

Beweis: • Preuve:

1. $X(t)$ ist die monoton wachsende absolute Häufigkeit ≤ 0
 - $X(t)$ est la fréquence absolue qui s'accroît de façon monotone ≤ 0
 - $\leadsto X(t + \Delta t) = X(t) + X(\Delta t)$
 - $\leadsto p_0(t + \Delta t) = P(X(t + \Delta t) = 0) = P(X(t) + X(\Delta t) = 0) = P(X(t) = 0 \wedge X(\Delta t) = 0)$
 - $=_{|\in A} P(X(t) = 0) \cdot P(X(\Delta t) = 0) = p_0(t) \cdot p_0(\Delta t) \approx p_0(t) \cdot (1 - p_1(\Delta t)) \approx p_0(t) \cdot (1 - a \cdot \Delta t)$
 - $= p_0(t) - p_0(t) \cdot a \cdot \Delta t \Rightarrow \frac{p_0(t + \Delta t) - p_0(t)}{\Delta t} = \frac{p_0(t) - p_0(t) \cdot a \cdot \Delta t - p_0(t)}{\Delta t} = -p_0(t) \cdot a$
 - $\Rightarrow \frac{dp_0(t)}{dt} = -p_0(t) \cdot a, \quad p_0(0) = \underbrace{P(X(0) = 0)}_{O.k.} = 0$
 - $\Rightarrow$ **AWP** • **PVI** $\leadsto$ bekannte Lösung: • *Solution connue:* $p_0(t) = e^{-at}$
2. $p_m(t + \Delta t) = P(X(t + \Delta t) = m) = P(X(t) + X(\Delta t) = m) =$
 $P((X(t) = m \wedge X(\Delta t) = 0) \vee (X(t) = m - 1 \wedge X(\Delta t) = 1) \vee \dots \vee (X(t) = 0 \wedge X(\Delta t) = m))$
 $=_{|\in A} P(X(t) = m) \cdot P(X(\Delta t) = 0) + P(X(t) = m - 1) \cdot P(X(\Delta t) = 1) + \dots + P(X(t) = 0) \cdot P(X(\Delta t) = m)$
 $= p_m(t) \cdot p_0(\Delta t) + p_{m-1}(t) \cdot p_1(\Delta t) + p_{m-2}(t) \cdot p_2(\Delta t) + \dots + p_0(t) \cdot p_m(\Delta t)$
 $\leq p_m(t) \cdot p_0(\Delta t) + p_{m-1}(t) \cdot p_1(\Delta t) + p_{m-2}(t) \cdot 1 + \dots + p_0(t) \cdot 1$
 $= p_m(t) \cdot p_0(\Delta t) + p_{m-1}(t) \cdot p_1(\Delta t) + \sum_{k=1}^m p_k(t) \approx p_m(t) \cdot p_0(\Delta t) + p_{m-1}(t) \cdot p_1(\Delta t) + \varphi(\Delta t) \cdot (\Delta t)^2,$
 $\Delta t \rightarrow 0 \Rightarrow \varphi(\Delta t) \cdot (\Delta t)^2 \rightarrow 0$
 $\leadsto p_m(t + \Delta t) \approx p_m(t) \cdot p_0(\Delta t) + p_{m-1}(t) \cdot p_1(\Delta t) = p_m(t) \cdot (1 - p_1(\Delta t)) + p_{m-1}(t) \cdot p_1(\Delta t)$
 $= p_m(t) \cdot (1 - a \cdot \Delta t) + p_{m-1}(t) \cdot a \cdot \Delta t = p_m(t) - p_m(t) \cdot a \cdot \Delta t + p_{m-1}(t) \cdot a \cdot \Delta t \Rightarrow$
 $\frac{p_m(t + \Delta t) - p_m(t)}{\Delta t} = \frac{p_m(t) - p_m(t) \cdot a \cdot \Delta t + p_{m-1}(t) \cdot a \cdot \Delta t - p_m(t)}{\Delta t} = -p_m(t) \cdot a + p_{m-1}(t) \cdot a$
 $\Rightarrow \frac{dp_m(t)}{dt} = -a \cdot p_m(t) + a \cdot p_{m-1}(t), \quad X(0) = 0 \leadsto P(X(0) = m > 0) = p_{m>0}(0) = 0$
 $\Rightarrow$ **AWP** • **PVI** $\leadsto$
 - (a) $m = 1 \leadsto \frac{dp_1(t)}{dt} = -a \cdot p_1(t) + a \cdot p_0(t) = -a \cdot p_1(t) + a \cdot e^{-at} = , \quad p_1(0) = 0$
 $\leadsto$ **AWP** • **PVI** $\leadsto$ bekannte Lösung: • *Solution connue:* $p_1(t) = a \cdot t \cdot e^{-at}$
 - (b) $m = 2 \leadsto \frac{dp_2(t)}{dt} = -a \cdot p_2(t) + a \cdot p_1(t) = -a \cdot p_2(t) + a \cdot a \cdot t \cdot e^{-at} = , \quad p_2(0) = 0$
 $\leadsto$ **AWP** • **PVI** $\leadsto$ bekannte Lösung: • *Solution connue:* $p_2(t) = \frac{(a \cdot t)^2}{2} \cdot e^{-at}$
 - (c) Induktion • *Induction* $\leadsto p_m(t) = \frac{(a \cdot t)^m}{m!} \cdot e^{-at}$

Konsequenz: • Conséquence:

Die Wahrscheinlichkeit, dass im Zeitintervall $[t_0, t_0 + \Delta t]$ kein Ereignis eintritt, ist unabhängig von t_0 und somit $= p_0(\Delta t) = e^{-a \Delta t}$.

• *La probabilité que dans l'intervalle temporel $[t_0, t_0 + \Delta t]$ aucun événement n'est réalisé est indépendante de t_0 et par conséquent $= p_0(\Delta t) = e^{-a \Delta t}$.*

Sei T = Wartezeit zwischen zwei Ereignissen, z.B. zwischen zwei Ausfällen von Bauteilen oder Maschinen.

• *Soit T = temps'attente entre deux événements, par exemple entre deux pannes (défaillances) de composants ou machines.*

$\leadsto R(t) := P(T > t) = p_0(t) = e^{-at}$ = Wahrscheinlichkeit, dass wegen $T > t$ bis t kein Bauteil ausfällt
 • = probabilité qu'à cause de $T > t$ jusqu'à t , aucun composant tombe en panne

$\leadsto F(t) := 1 - R(t)$ ist die **Überlebenswahrscheinlichkeit** • *est la probabilité de survivre*

$\leadsto F(t) := 1 - R(t) = P(T \leq t) = 1 - e^{-at} \leadsto$ Exponentialverteilung • *Répartition exponentielle*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Wie im letzten Satz • *Comme au théorème dernier*
 $R(t) := P(T > t)$

Beh.: • **Thè.:**

$F(t) := P(T \leq t) = 1 - e^{-at}$

Konsequenz: • **Conséquence:** Hier war die Häufigkeit $X(t)$ poissonverteilt. Daraus folgt somit hier, dass die Wartezeiten T nach $P(T \leq t) = 1 - e^{-at}$ exponentialverteilt sind.

• *Ici, la fréquence $X(t)$ était distribuée selon la loi de Poisson. Par conséquent on peut donc conclure que les attentes T sont réparties selon la loi exponentielle $P(T \leq t) = 1 - e^{-at}$.*

Bemerkung: • **Remarque:** Von diesem Sachverhalt kann man auch die Umkehrung beweisen.
 • *On peut aussi prouver le théorème inverse.*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$F(t) = P(T \leq t) = 1 - e^{-at}$, $\lambda = a \cdot t$

Beh.: • **Thè.:**

$p_m(t) = \frac{\lambda \cdot t}{m!} \cdot e^{-\lambda \cdot t}$

5.4.2 Warteschlangemodelle — Modèles de files d'attente

In den praktischen Anwendungen der Mathematik im Bereich der Planung spielt das Problem der Warteschlangen eine grössere Rolle (Kundenschalter, Drucker, Übermittlung u.s.w.). Betrachten wir z.B. das Problem von wartenden Kunden an einer Kasse. Die Anzahl der Kunden, die sich pro Zeiteinheit vor der Kasse einfinden, ist meist Poissonverteilt. Die Zeit aber, die pro Kunde zur Abfertigung benötigt wird, ist meist exponentialverteilt. Beide Prozesse spielen hier also zusammen. Die Kombination bestimmt ein Warteschlangensystem. Die weiterreichende Behandlung des Problems ist wieder Gegenstand einer ausgedehnten Theorie, die den Rahmen dieses Kurses sprengt.

• *Lors de l'usage pratique des mathématiques dans le domaine de la planification, le problème des queues ou files d'attente joue un grand rôle (quiches, imprimante, transmission, etc.). Considérons par ex. le problème de clients qui font la queue devant une caisse. Le nombre de clients qui se trouve devant la caisse par unité temporelle est distribué d'après Poisson. Le temps que le client emploie pour être servi est pour la plupart distribué exponentiellement. Les deux processus entrent donc en relation. Leur combinaison détermine un système des files d'attente. Pour approfondir ce problème intensivement, il nous faudrait une ample théorie, ce qui déborderait le cadre de ce cours.*

5.5 Schätzungen — Estimations

5.5.1 Schätzer, Punktschätzer — Estimations ponctuelles

Problem: • **Problème:**

Gegeben sei eine Stichprobe von Werten $\{x_1, x_2, \dots, x_n\}$ aus einer weiter nicht bekannten Grundgesamtheit U . Die Werte werden von einer Variablen X angenommen, für deren Verteilungsfunktion ein Modell vorliegen kann. Für U sollen Parameter wie μ, σ^2 u.s.w. mit Hilfe von bekannten empirischen Parametern der Stichprobe geschätzt werden. Als Methode verwenden wir die **Momentenmethode**. Die zu schätzenden Parameter sollen durch die Momente der Verteilung $\mu = E(X)$, $\sigma = E((X - \mu)^2)$ u.s.w. ausgedrückt werden. Dazu verwenden wir **Schätzfunktionen**, **Punktschätzer** oder **Schätzer**. **Schätzwerte** sind zahlenmässige Realisierungen von Punktschätzern. Möchte man hingegen Aussagen gewinnen über die Genauigkeit oder die Sicherheit von Schätzwerten, so kann man dazu **Schätzintervalle** oder **Konfidenzintervalle (Vertrauensintervalle)** berechnen.

• *Soit donné un échantillon de valeurs $\{x_1, x_2, \dots, x_n\}$ d'un ensemble de valeurs fondamentale U qui soit inconnu dans sa totalité. Les valeurs sont atteintes une variable X dont on peut avoir un modèle pour la fonction de répartition. Pour U il faudrait estimer des paramètres comme μ, σ^2 etc. à l'aide de paramètres empiriques et connus de l'échantillon. Comme méthode, nous utilisons la méthode des moments. Les paramètres à estimer devraient être exprimés par les moments de la répartition $\mu = E(X)$, $\sigma = E((X - \mu)^2)$ etc.. Dans cette intention nous utilisons des **fonctions d'estimation** ou des **estimations ponctuelles**. Les **valeurs estimées** sont des réalisations d'estimations ponctuelles en nombres. Si on veut par contre obtenir des connaissances sur l'exactitude ou la sécurité de valeurs d'estimation, on peut calculer des **intervalles estimés** ou des **intervalles de confiance**.*

Definition: • **Définition:**

Schätzfunktionen, **Punktschätzer** oder **Schätzer** sind Stichprobenfunktionen, durch die wir unbekannte theoretische Wahrscheinlichkeitsfunktionen oder Parameter einer Grundgesamtheit zufriedenstellend ersetzen wollen.

• *Des **fonctions d'estimation** ou des **valeurs estimées** sont des fonctions de l'échantillon par lesquelles nous voulons remplacer de façon satisfaisante des fonctions aléatoires théoriques ou des paramètres de l'ensemble fondamental.*

Trick: • **Truc:** Für die Stichprobenwerte x_i wollen wir jetzt individuell neue Variablen X_i einführen:

• *Pour les valeurs x_i de l'échantillon nous voulons maintenant introduire individuellement de nouvelles variables X_i :*

Annahme: x_1 sei ein Wert für X_1 , x_2 ein Wert für X_2 u.s.w.. Die so postulierten Variablen haben alle dieselbe Verteilung wie X .

• **Hypothèse:** x_1 soit une valeur pour X_1 , x_2 une valeur pour X_2 etc.. Les variables ainsi postulées ont toutes la même distribution que X .

Bsp.: • **Exemple:**

Schätzer für μ : • *Valeur estimée pour μ :* $\bar{X} = \bar{X}(n) = \frac{1}{n} \sum_{i=1}^n X_i$

Schätzer für σ^2 : • *Valeur estimée pour σ^2 :* $S^2 = S^2(n) = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ oder • *ou*

$$Q^2 = Q^2(n) = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{(n-1)S^2}{n}$$

Problem: • **Problème:**

Wie gut oder tauglich sind solche Schätzer? • *Quelle est la qualité de tels estimations ponctuelles?*

5.5.2 Erwartungstreue — Propriété d'être sans biais**Definition:** • **Définition:**

Ein Schätzer Y resp. Y_n für einen Parameter Υ der Grundgesamtheit heisst **erwartungstreu** oder **unverzerrt**, wenn der Erwartungswert von Y mit dem Erwartungswert von Υ übereinstimmt. D.h. wenn gilt:

• *Une valeur estimée Y resp. Y_n pour un paramètre Υ de l'ensemble de base s'appelle **non biaisé**, si la valeur probable de Y est la même que la valeur probable de Υ . C.v.d. s'il vaut:*

$$E(Y) = E(\Upsilon)$$

Bemerkung: • **Remarque:**

Gegeben sei eine Stichprobe $\{x_1, x_2, \dots, x_n\}$. Seien die Variablen $X_1, X_2, \dots, X_n$ wieder verteilt wie X . Speziell gilt dann $E(X_i) = E(X)$, $i = 1, \dots, n$, sowie $\sigma^2 = E((X - \mu)^2) = E((X_i - \mu)^2)$, $i = 1, \dots, n$. Der Wertevektor $(x_1, x_2, \dots, x_n)^T$ gibt dann eine mögliche Kombination von Werten, die vom Variablenvektor $(X_1, X_2, \dots, X_n)^T$ angenommen werden können (T bedeutet „transponiert“). Dann sind der Stichprobenmittelwert und die empirische Stichprobenvarianz erwartungstreue Schätzer für μ und σ^2 .

• *Soit donné un échantillonnage $\{x_1, x_2, \dots, x_n\}$. Les variables $X_1, X_2, \dots, X_n$ soient encore réparties comme X . Il vaut donc spécialement $E(X_i) = E(X)$, $i = 1, \dots, n$ ainsi que $\sigma^2 = E((X - \mu)^2) = E((X_i - \mu)^2)$, $i = 1, \dots, n$. Alors le vecteur de valeurs $(x_1, x_2, \dots, x_n)^T$ donne une combinaison possible de valeurs qui sont atteintes par le vecteur de variables $(X_1, X_2, \dots, X_n)^T$ (T signifie „transposé“). Alors la moyenne de l'échantillonnage et la variance empirique sont des valeurs estimées pour μ et σ^2 non biaisées.*

Symbol: • **Symbole:** $X_i \sim X \rightsquigarrow X_i$ verteilt wie X • X_i réparti comme X

Satz: • **Théorème:** (Mittelwert erwartungstreu)

• *(Moyenne fidèle à l'espérance mathématique)*

Vor.: • **Hyp.:**

$$X_1, X_2, \dots, X_n \sim X$$

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i, \quad \mu = E(X)$$

$\bar{X}$ Schätzer für μ • $\bar{X}$ valeur estimée pour μ

Beh.: • **Thè.:**

$$E(\bar{X}) = \mu = E(X)$$

Beweis: • **Preuve:**

$$E(\bar{X}) = E\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n} \sum_{i=1}^n E(X_i) = \frac{1}{n} \sum_{i=1}^n \mu = \frac{1}{n} \cdot n \cdot \mu = \mu = E(X)$$



Satz: • Théorème:

(Empirische Varianz erwartungstreu

• *(Variance empirique fidèle à l'espérance mathématique)***Vor.: • Hyp.:**

$$X_1, X_2, \dots, X_n \sim X$$

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2, \quad \sigma^2 = E((X - \mu)^2)$$

 S^2 Schätzer für σ^2 • S^2 valeur estimée pour σ^2 **Beh.: • Thè.:**

$$E(S^2) = \sigma^2 = E((X - \mu)^2)$$

Beweis: • Preuve:

$$\begin{aligned} E(S^2) &= E\left(\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2\right) = \frac{1}{n-1} \sum_{i=1}^n E((X_i - \bar{X})^2) = \frac{1}{n-1} \sum_{i=1}^n E(((X_i - \mu) + (\mu - \bar{X}))^2) \\ &= \frac{1}{n-1} \sum_{i=1}^n E((X_i - \mu)^2 + 2(X_i - \mu)(\mu - \bar{X}) + (\mu - \bar{X})^2) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n E((X_i - \mu)^2) + E(2(X_i - \mu)(\mu - \bar{X})) + E((\mu - \bar{X})^2) \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 + 2E((X_i - \mu)\left(\frac{n}{n}\mu - \frac{1}{n} \sum_{j=1}^n X_j\right)) + E\left(\left(\frac{n}{n}\mu - \frac{1}{n} \sum_{j=1}^n X_j\right)^2\right) \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 + \frac{2}{n^2} E((X_i - \mu)(n\mu - \sum_{j=1}^n X_j)) + \frac{1}{n^2} E((n\mu - \sum_{j=1}^n X_j)^2) \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 + \frac{2}{n^2} E((X_i - \mu) \sum_{j=1}^n (\mu - X_j)) + \frac{1}{n^2} E\left(\left(\sum_{j=1}^n (\mu - X_j)\right)^2\right) \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 + \frac{2}{n^2} \sum_{j=1}^n E((X_i - \mu)(\mu - X_j)) + \frac{1}{n^2} \sum_{j=1}^n E((\mu - X_j)^2) \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 + \frac{2}{n^2} \sum_{j=1}^n E(\underbrace{-(X_i - \mu)(X_j - \mu)}_{i \neq j: \text{unabh.} \bullet \text{ indep.} \Rightarrow = 0}) + \frac{n}{n^2} \sigma^2 \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 - \frac{2}{n^2} n E((X_i - \mu)^2) + \frac{n}{n^2} \sigma^2 \right) = \frac{1}{n-1} \left(\sum_{i=1}^n \sigma^2 - \frac{2}{n} \sigma^2 + \frac{1}{n} \sigma^2 \right) = \frac{\sigma^2}{n-1} \left(n - \frac{2n}{n} + \frac{n}{n} \right) \\ &= \frac{\sigma^2}{n-1} (n - 2 + 1) = \frac{\sigma^2}{n-1} (n - 1) = \sigma^2 \quad \text{☺} \end{aligned}$$

Es gilt: • *Il vaut:*

$$Q^2 = Q^2(n) = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{(n-1)S^2}{n} \Rightarrow E(Q^2) = \frac{(n-1)}{n} E(S^2) = \frac{(n-1)}{n} \sigma^2 \neq \sigma^2 \rightsquigarrow$$

Korollar: • Corollaire: Q^2 nicht erwartungstreu. • Q^2 non biaisé.**Bemerkung: • Remarque:**Weiter kann man zeigen, dass S selbst nicht allgemein erwartungstreu ist (vgl. Lit. Storm, Bibl. A12).• *En outre on peut démontrer que S n'est pas non biaisé (voir lit. Storm, Bibl. A12).*

Bemerkung: • **Remarque:**

Eine schwächere Bedingung als die Erwartungstreue ist die asymptotische Erwartungstreue, für die man fordert:

- *Une condition plus faible que la fidélité à l'espérance est la fidélité à l'espérance asymptotique:*

$$\lim_{n \rightarrow \infty} E(Y_n) = E(\Upsilon)$$

5.5.3 Konsistenz — Consistance

Sei wieder $X_i \sim X$. • *Soit de nouveau $X_i \sim X$.*

Definition: • **Définition:**

Ein Schätzer Y resp. Y_n für einen Parameter Υ der Grundgesamtheit heisst **konsistent** oder **schwach konsistent** oder **stichhaltig**, wenn gilt:

- *Une valeur estimée Y resp. Y_n pour un paramètre Υ de l'ensemble de base s'appelle **consistant** ou **consistant de façon faible**, s'il vaut:*

$$\forall \varepsilon > 0 \quad \lim_{n \rightarrow \infty} P(|Y_n - \Upsilon| < \varepsilon) = 1$$

D.h. mit wachsendem Stichprobenumfang konvergiert die Wahrscheinlichkeit, dass sich der Schätzer beliebig wenig vom zu schätzenden Parameter unterscheidet, gegen 1.

- *C.-à.-d. avec la taille d'échantillon grandissante la probabilité que l'estimation ponctuelle se distingue aussi peu que possible du paramètre à estimer, converge vers 1.*

Satz: • **Théorème:**

Vor.: • **Hyp.:**

$$\begin{aligned} \bar{X} &= \bar{X}(n) = \frac{1}{n} \sum_{i=1}^n X_i \\ S^2 &= S^2(n) = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \\ Q^2 &= Q^2(n) = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{(n-1) S^2}{n} \end{aligned}$$

Beh.: • **Thè.:**

1. $\bar{X}$ ist ein konsistenter Schätzer für μ
 - $\bar{X}$ est une estimation consistante pour μ
2. $\lim_{n \rightarrow \infty} \text{Var}(S^2) \rightarrow 0$, $\lim_{n \rightarrow \infty} \text{Var}(Q^2) \rightarrow 0$ oder • *ou $X_i \in [-M, M]$*
 $\Rightarrow S^2$ und Q^2 sind konsistente Schätzer für σ^2
 - S^2 et Q^2 sont des estimations consistantes pour σ^2

Zum Beweis verwenden wir das Lemma:

- *Pour la preuve nous utilisons le lemme:*

Lemma: • **Lemme:****Vor.:** • **Hyp.:**

$$\begin{aligned}\forall_i X_i &\sim X \\ E(X_i) &= E(X) = \mu \\ \text{Var}(X_i) &= \text{Var}(X) = \sigma^2\end{aligned}$$

Beh.: • **Thè.:**

$$1. E(\bar{X}) = \mu$$

$$2. \text{Var}(\bar{X}) = E((\bar{X} - \mu)^2) = E(\bar{X}^2) - \mu^2 = \frac{\sigma^2}{n}$$

Beweis: • **Preuve:**

(Lemma) • (Lemme)

$$\begin{aligned}1. E(\bar{X}) &= E\left(\frac{1}{n} \sum_{i=1}^n X_i\right) = \frac{1}{n} \sum_{i=1}^n E(X_i) = \frac{1}{n} \sum_{i=1}^n \mu = \frac{n}{n} \mu = \mu \\ 2. \text{Var}(\bar{X}) &= E((\bar{X} - \mu)^2) = E\left(\left(\frac{1}{n} \sum_{i=1}^n X_i - \frac{n}{n} \mu\right)^2\right) = E\left(\frac{1}{n^2} \left(\sum_{i=1}^n X_i - (n\mu)\right)^2\right) \\ &= \frac{1}{n^2} E\left(\left(\sum_{i=1}^n (X_i - \mu)\right)^2\right) = \frac{1}{n^2} E\left(\sum_{i,j=1}^n (X_i - \mu)(X_j - \mu)\right) = \frac{1}{n^2} \sum_{i,j=1}^n E\left(\underbrace{(X_i - \mu)(X_j - \mu)}_{i \neq j: \text{unabh.} \bullet \text{indép.} \Rightarrow =0}\right) \\ &= \frac{1}{n^2} \sum_{i=1}^n E((X_i - \mu)^2) = \frac{1}{n^2} n \sigma^2 = \frac{\sigma^2}{n}\end{aligned}$$

Zum Beweis des Satzes: • **Quant à la preuve du théorème:**

Verwende Tschebyscheff: • *Utiliser Tschebyscheff:* $P(|Y - E(Y)| \geq \varepsilon) \leq \frac{\text{Var}(Y)}{\varepsilon^2}$

$$\begin{aligned}1. \lim_{n \rightarrow \infty} P(|Y_n - \Upsilon| < \varepsilon) &= \lim_{n \rightarrow \infty} P(|\bar{X} - \mu| < \varepsilon) = \lim_{n \rightarrow \infty} (1 - P(|\bar{X} - \mu| \geq \varepsilon)) \\ &= 1 - \lim_{n \rightarrow \infty} P(|\bar{X} - \mu| \geq \varepsilon), \quad P(|\bar{X} - \mu| \geq \varepsilon) \leq \frac{\text{Var}(\bar{X})}{\varepsilon^2} = \frac{\sigma^2}{n \varepsilon^2} \rightarrow 0, \quad n \rightarrow \infty \\ &\Rightarrow \lim_{n \rightarrow \infty} P(|\bar{X} - \mu| \geq \varepsilon) = 0 \Rightarrow \lim_{n \rightarrow \infty} P(|\bar{X} - \mu| < \varepsilon) = 1 \\ 2. \lim_{n \rightarrow \infty} P(|Y_n - \Upsilon| < \varepsilon) &= \lim_{n \rightarrow \infty} P(|S^2 - \sigma^2| < \varepsilon) = 1 - \lim_{n \rightarrow \infty} P(|S^2 - \sigma^2| \geq \varepsilon), \\ P(|S^2 - \sigma^2| \geq \varepsilon) &= P(|S^2 - E(S^2)| \geq \varepsilon) \leq \frac{\text{Var}(S^2)}{\varepsilon^2} = \frac{1}{\varepsilon^2} \text{Var}\left(\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2\right) \\ (X_i - \bar{X}), (X_j - \bar{X}) &\text{unabhängig für } i \neq j \bullet (X_i - \bar{X}), (X_j - \bar{X}) \text{ indépendants pour } i \neq j \\ (\text{p. 137}) \rightsquigarrow \frac{1}{\varepsilon^2} \text{Var}\left(\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2\right) &= \frac{1}{\varepsilon^2} \sum_{i=1}^n \text{Var}\left(\frac{1}{n-1} (X_i - \bar{X})^2\right) \\ &= \frac{1}{\varepsilon^2} \sum_{i=1}^n E\left(\left(\frac{1}{n-1} (X_i - \bar{X})^2 - \frac{1}{n-1} E((X_i - \bar{X})^2)\right)^2\right) \\ &= \frac{1}{\varepsilon^2 (n-1)^2} \sum_{i=1}^n \underbrace{E(((X_i - \bar{X})^2 - E((X_i - \bar{X})^2))^2)}_{\leq \text{Const} = M)} \leq \frac{M n}{\varepsilon^2 (n-1)^2} \rightarrow 0, \quad n \rightarrow \infty\end{aligned}$$

Weitere Begriffe — Autres notions

Wirksamkeit, Effizienz, Suffizienz, Wirkungsgrad etc. vgl. Lit..

- *Efficacité, exhaustivité, rendement etc, voir lit.*

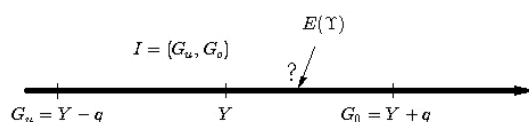
5.5.4 Vertrauensintervalle I — Intervalles de confiance I**Begriffe — Notions**

Möchte man Aussagen gewinnen über die Genauigkeit oder die Sicherheit von Schätzwerten, so kann man dazu **Schätzintervalle** oder **Konfidenzintervalle (Vertrauensintervalle)** berechnen.

- *Si on veut gagner des connaissances sur l'exactitude ou la sécurité de valeurs d'estimation, on peut calculer des intervalles estimés ou des intervalles de confiance.*

Sei Y resp. Y_n ein Schätzer für einen Parameter Υ der Grundgesamtheit.

- *Soit Y resp. Y_n une estimation ponctuelle pour un paramètre Υ de l'ensemble de base.*



Sei • *Soit* $I = (G_u, G_o)$

$$P(G_u < E(\Upsilon) < G_o) = \varepsilon = 1 - \alpha$$

Bsp.: • **Exemple:**

$$\varepsilon = 0.95 \hat{=} 95\% \text{ oder } \bullet \text{ ou } \varepsilon = 0.99 \hat{=} 99\%$$

Definition: • **Définition:**

$I = (G_u, G_o)$ heisst **Vertrauensintervall** oder **Konfidenzintervall** zur **Vertrauenswahrscheinlichkeit** $\varepsilon = 1 - \alpha$.

- $I = (G_u, G_o)$ s'appelle **intervalle de confiance** à la probabilité de confiance $\varepsilon = 1 - \alpha$.

Bemerkung: • **Remarque:**

Vertrauensintervalle interessieren in der Praxis!

- *Les intervalles de confiance intéressent en pratique!*

Beispiel Mittelwert — Exemple moyenne

Bsp.: • **Exemple:**

Sei X gaussverteilt mit schon bekannter Varianz σ^2 . Gesucht ist das Konfidenzintervall für μ zu $\varepsilon = 1 - \alpha$.

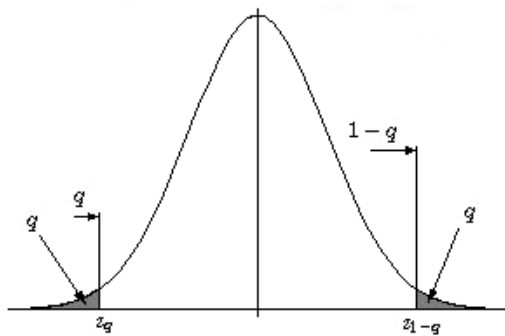
- *Soit X réparti d'après Gauss avec la variance σ^2 qu'on connaît. On cherche l'intervalle de confiance μ pour $\varepsilon = 1 - \alpha$.*

Wir wissen: • *Nous savons:*

$$X \in N(\mu, \sigma^2) \Rightarrow \bar{X} \in N(\mu, \frac{\sigma^2}{n})$$

$$X \in N(\mu, \sigma^2) \Leftrightarrow Z = \frac{X - \mu}{\sigma} \in N(0, 1)$$

$$\bar{X} \in N(\mu, \frac{\sigma^2}{n}) \Leftrightarrow \bar{Z} = \frac{\bar{X} - \mu}{(\sigma/\sqrt{n})} \in N(0, 1)$$



Für das Weitere studieren wir die normierte Normalverteilung $N(0, 1)$. Wenn man mit Rechnern arbeitet, ist die Einschränkung auf die Standardsituation nicht notwendig. Wenn man wie früher mit Tabellen arbeitet aber schon.)

• *Pour ce qu'on va faire, nous étudions la répartition normale $N(0, 1)$. Si on travaille en utilisant des calculateurs, la restriction à la situation standard n'est pas nécessaire. Si on travaille comme jadis avec les tableaux, il faut la restriction.*

Bekannt: • *Nous savons:*

$$1. F(x) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x e^{-\frac{(t-\mu)^2}{2\sigma^2}} dt = P(X \leq x) = \Phi(x; \mu, \sigma^2)$$

$$2. \Phi(z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt = P(Z \leq z) = \Phi(z; 0, 1)$$

$$3. Z = \frac{X - \mu}{\sigma} \in N(0; 1) \Leftrightarrow X \in N(\mu; \sigma^2)$$

$$\leadsto \text{Sei } \bullet \text{ Soit } \Phi(z) = \Phi(z; 0, 1) = P(X \leq z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt \leadsto \Phi(b) = P(X \leq b)$$

$$\leadsto \text{Sei } \bullet \text{ Soit } q := \Phi(z_q) = \Phi(z_q; 0, 1) \leadsto q\text{-Quantil (Ordnung } q) \bullet \text{ quantile } q \text{ (Ordre } q)$$

Man sieht aus der Skizze: • *On voit dans l'esquisse:*

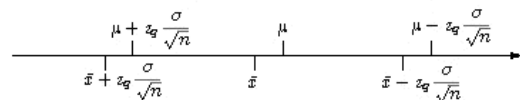
$$\text{Formel: } \bullet \text{ Formule: } z_q = -z_{1-q}, \quad \Phi(z_{1-q}) = 1 - q = 1 - \Phi(z_q)$$

$$\text{Sei } \bullet \text{ Soit } z \leadsto N(0, 1), \quad \bar{X} \leadsto N(\mu, \frac{\sigma^2}{n}) \text{ und } \bullet \text{ et } F(\bar{X}) = F(\bar{X}, \mu, \frac{\sigma^2}{n}) = \Phi(\frac{\bar{X} - \mu}{\sigma/\sqrt{n}}; 0, 1) = q$$

$$\leadsto \text{Transformation: } \bullet \text{ Transformation: } Z = \frac{\bar{X} - \mu}{\sigma/\sqrt{n}}, \quad \bar{X} = \mu + Z \cdot \frac{\sigma}{\sqrt{n}}$$

$$\text{Sei } \bullet \text{ Soit } \bar{x} = \frac{1}{n} \sum_{i=1}^n x_i \text{ (Stichprobenmittelwert)} \bullet \text{ (Valeur moyenne de l'échantillonnage)}$$

$$\text{Sei } \bullet \text{ Soit } z = z_q$$



$$\leadsto \text{Es gilt: } \bullet \text{ Il vaut: } (z_q < 0 !!!)$$

$$\begin{aligned} \mu \in [\bar{x} + z_q \frac{\sigma}{\sqrt{n}}, \bar{x} - z_q \frac{\sigma}{\sqrt{n}}] &\Leftrightarrow (\mu - \bar{x}) \in [z_q \frac{\sigma}{\sqrt{n}}, -z_q \frac{\sigma}{\sqrt{n}}] \Leftrightarrow (\bar{x} - \mu) \in [z_q \frac{\sigma}{\sqrt{n}}, -z_q \frac{\sigma}{\sqrt{n}}] \\ &\Leftrightarrow |\bar{x} - \mu| \leq |z_q| \frac{\sigma}{\sqrt{n}} \Leftrightarrow \bar{x} \in [\mu + z_q \frac{\sigma}{\sqrt{n}}, \mu - z_q \frac{\sigma}{\sqrt{n}}] \end{aligned}$$

$$\text{Wichtig: } \bullet \text{ Important: } \leadsto \mu \in [\bar{x} + z_q \frac{\sigma}{\sqrt{n}}, \bar{x} - z_q \frac{\sigma}{\sqrt{n}}] \Leftrightarrow \bar{x} \in [\mu + z_q \frac{\sigma}{\sqrt{n}}, \mu - z_q \frac{\sigma}{\sqrt{n}}]$$

$$\begin{aligned}
\leadsto \varepsilon = 1 - \alpha &= P(\mu \in [\bar{X} + z_q \frac{\sigma}{\sqrt{n}}, \bar{X} - z_q \frac{\sigma}{\sqrt{n}}]) = P(\bar{X} \in [\mu + z_q \frac{\sigma}{\sqrt{n}}, \mu - z_q \frac{\sigma}{\sqrt{n}}]) \\
&= P(\mu + z_q \frac{\sigma}{\sqrt{n}} \leq \bar{X} \leq \mu - z_q \frac{\sigma}{\sqrt{n}}) = F(\mu - z_q \frac{\sigma}{\sqrt{n}}) - F(\mu + z_q \frac{\sigma}{\sqrt{n}}) = (1 - q) - q = \Phi(z_{1-q}) - \Phi(z_q) \\
&= \Phi(z_{1-q}) - (1 - \Phi(z_{1-q})) = 2\Phi(z_{1-q}) - 1
\end{aligned}$$

$$\Rightarrow 2 - \alpha = 2\Phi(z_{1-q}), \quad \Phi(z_{1-q}) = 1 - \frac{\alpha}{2} = \int_{-\infty}^{z_{1-q}} \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt$$

Formel: • **Formule:** $\Phi(z_{1-q}) = 1 - \frac{\alpha}{2} = \int_{-\infty}^{z_{1-q}} \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt, \quad \Phi(z_q) = \frac{\alpha}{2} = \int_{-\infty}^{z_q} \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt$

z_q nennen wir daher **zweiteilige α -Schranke** der normierten Gaussverteilung Φ . z_q ist das $\frac{\alpha}{2}$ -Quantil.

• *Par conséquent nous appelons z_q borne α bipartite de la répartition de Gauss normalisée Φ . z_q est le quantile de la valeur $\frac{\alpha}{2}$.*

Zahlenbeispiel: • *Exemple avec des nombres:* $\alpha = 0.05 \hat{=} 5\%$

$$\leadsto \Phi(z_{1-q}) = \int_{-\infty}^{z_{1-q}} \frac{1}{\sqrt{2\pi}} e^{-\frac{t^2}{2}} dt = 1 - \frac{0.05}{2} = 0.975$$

Mathematica-Programm: • **Programme en Mathematica:**

```
| FindRoot[Integrate[E^(-t^2/2)/Sqrt[2 Pi], {t, -Infinity, z}] == 0.975, {z, 1}]|
```

Output: • **Output:**

```
| {z -> 1.95996}|
```

$\leadsto$ Wir finden daher $z_{1-q} \approx 1.960$. • *Nous trouvons donc $z_{1-q} \approx 1.960$.*

Seien • *Soient* $\bar{x} = 156.38, \sigma^2 = 0.20, n = 16$.

$$\begin{aligned}
\leadsto \varepsilon = 1 - \alpha = 0.95 &= P(\mu \in [\bar{X} + z_q \frac{\sigma}{\sqrt{n}}, \bar{X} - z_q \frac{\sigma}{\sqrt{n}}]) \\
&\approx P(\mu \in [156.38 - 1.960 \frac{\sqrt{0.20}}{\sqrt{16}}, 156.38 + 1.960 \frac{\sqrt{0.20}}{\sqrt{16}}]) \\
&\approx P(\mu \in [156.161, 156.599]) \leadsto I = [156.161, 156.599] = [a, b], \quad b - a \approx 0.438269
\end{aligned}$$

I ist das Vertrauensintervall zur Vertrauenswanrscheinlichkeit $\varepsilon = 1 - \alpha = 0.95$. Die Wahrscheinlichkeit, dass μ ausserhalb dieses Intervalls liegt, ist $\alpha = 0.05$.

• *I est l'intervalle de confiance pour la probabilité de confiance $\varepsilon = 1 - \alpha = 0.95$. La probabilité que μ est situé à l'extérieur de cet intervalle, est $\alpha = 0.05$.*

5.5.5 Wichtige Testfunktionen — Fonctions de test importantes

Oft kennt man μ und σ nicht. Das bedeutet, dass man mit Schätzungen für diese Grössen arbeiten muss. Es stellt sich dann die Frage nach den Typen der entsprechendnen Verteilungsfunktionen, falls man μ

und σ durch Schätzungen ersetzt. Da die Resultate eigentlich erstaunlich sind, müssen wir dazu einige Tatsachen herleiten.

• *Souvent on ne connaît ni μ ni σ . Ça signifie qu'on doit travailler avec des estimations. Alors il se pose la question d'après les types des fonctions de répartition correspondantes, si on doit remplacer μ et σ par des estimations. Comme les résultats sont proprement étonnants, nous devons maintenant déduire quelques faits.*

$$\text{Sei } \bullet \text{ Soit } (\bar{X} = \sum_{i=1}^n X_i \Rightarrow \sum_{i=1}^n (X_i - \bar{X}) = (\sum_{i=1}^n X_i) - n\bar{X} = \sum_{i=1}^n X_i - n \frac{1}{n} \sum_{i=1}^n X_i = \sum_{i=1}^n X_i - \sum_{i=1}^n X_i = 0$$

$$U_i = (X_i - \bar{X}), \quad \sum_{i=1}^n U_i = 0 \leadsto \text{Die Variablen } U_i = (X_i - \bar{X}) \text{ sind abhängig.}$$

• *Les variables $U_i = (X_i - \bar{X})$ sont dépendantes.*

$$\leadsto \text{Z.B. } \bullet \text{ P.ex. } U_1 = - \sum_{i=2}^n U_i$$

$$\leadsto 0 = 0^2 = (\sum_{i=1}^n U_i)^2 = \sum_{i=1}^n U_i^2 + \sum_{(i \neq k \wedge i, k=1)}^n U_i \cdot U_k$$

$$\Rightarrow U_1^2 = - \sum_{i=2}^n U_i^2 + \sum_{(i \neq k \wedge i, k=1)}^n U_i \cdot U_k = - \sum_{i=2}^n U_i^2 + \sum_{(i \neq k \wedge i, k=2)}^n U_i \cdot U_k + U_1 \cdot \sum_{k=2}^n U_k + U_1 \cdot \sum_{i=2}^n U_i$$

$$= - \sum_{i=2}^n U_i^2 + \sum_{(i \neq k \wedge i, k=2)}^n U_i \cdot U_k + (- \sum_{i=2}^n U_i) \cdot \sum_{k=2}^n U_k + (- \sum_{k=2}^n U_k) \cdot \sum_{i=2}^n U_i = h(U_2, U_3, \dots, U_n)$$

$$\leadsto U_1^2 = h(U_2, U_3, \dots, U_n)$$

$\leadsto$ Die Variablen $U_i^2 = (X_i - \bar{X})^2$ sind abhängig. • *Les variables $U_i^2 = (X_i - \bar{X})^2$ sont dépendantes.*

Lemma: • **Lemme:** Die Variablen $U_i = (X_i - \bar{X})$ und auch die Variablen $U_i^2 = (X_i - \bar{X})^2$ sind abhängig

• *Les variables $U_i = (X_i - \bar{X})$ et aussi les variables $U_i^2 = (X_i - \bar{X})^2$ sont dépendantes.*

Weiter ist unmittelbar klar, dass eine Koordinatentransformation $X_i \mapsto X_i - \mu$ resp. $\bar{X}_i \mapsto \bar{X}_i - \mu$ keinen Einfluss auf die Unabhängigkeit oder Abhängigkeit hat. Wir verwenden folgendes Symbol:

• *En outre il est évident qu'une transformation de coordonnées $X_i \mapsto X_i - \mu$ resp. $\bar{X}_i \mapsto \bar{X}_i - \mu$ n'a pas d'influence sur la dépendance ou indépendance. Nous utilisons le symbole suivant:*

Symbol: • **Symbole:** X_i unabhängig • X_i indépendant $\Leftrightarrow X_i \in \{\heartsuit\}$
 X_i abhängig • X_i dépendant $\Leftrightarrow X_i \in \{\spadesuit\}$

Lemma: • **Lemme:** **Vor.:** • **Hyp.:**

$$\{X_i\} \in \{\heartsuit\} \text{ resp. } \bullet \text{ resp. } \{X_i\} \in \{\spadesuit\}$$

Beh.: • **Thè.:**

$$\{X_i - \mu\} \in \{\heartsuit\} \text{ resp. } \bullet \text{ resp. } \{X_i - \mu\} \in \{\spadesuit\}$$

$$\leadsto X_i \in N(\mu, \sigma^2) \Rightarrow (X_i - \mu) \in N(0, \sigma^2)$$

$\leadsto X_i \in N(0, \sigma^2)$ bringt keine Einschränkung weiterer Bedingungen.

• $X_i \in N(0, \sigma^2)$ ne restreint pas d'autres conditions.

Sei • Soit $X_i \in N(0, \sigma^2)$, $Q_n = \sum_{i=1}^n (X_i - \bar{X})^2$

$\leadsto$ Q_n kann als Summe von $n - 1$ unabhängigen Quadraten geschrieben werden. Das wollen wir jetzt herleiten:

• Il est possible d'écrire Q_n comme somme de $n - 1$ variables au carré qui sont indépendantes. C'est ce que nous allons déduire:

Lemma: • **Lemme:** **Vor.:** • **Hyp.:**

$$X_i \in N(0, \sigma^2), \quad Q_n = \sum_{i=2}^n (X_i - \bar{X})^2$$

Beh.: • **Thè.:**

$$\exists_{Y_1, Y_2, \dots, Y_{n-1} \in N(0, \sigma^2), \in \{\heartsuit\}} : Q_n = \sum_{i=2}^{n-1} Y_i^2$$

Beweis: • **Preuve:**

(Lemma) • (Lemme) Sei

$$\begin{aligned} Y_1 &= \frac{1}{\sqrt{1 \cdot 2}} (X_1 - X_2) \\ Y_2 &= \frac{1}{\sqrt{2 \cdot 3}} (X_1 + X_2 - 2 X_3) \\ &\vdots \\ Y_{n-1} &= \frac{1}{\sqrt{(n-1) \cdot n}} (X_1 + X_2 + \dots + X_{n-1} - (n-1) X_n) \\ Y_n &= \frac{1}{\sqrt{n}} (X_1 + X_2 + \dots + X_{n-1} + X_n) \end{aligned}$$

$$\leadsto \begin{pmatrix} Y_1 \\ \vdots \\ Y_n \end{pmatrix} = \begin{pmatrix} a_{11} & \dots & a_{1n} \\ \vdots & & \vdots \\ a_{n1} & \dots & a_{nn} \end{pmatrix} \cdot \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix} = M \cdot \begin{pmatrix} X_1 \\ \vdots \\ X_n \end{pmatrix}$$

$$M = \begin{pmatrix} \frac{1}{\sqrt{1 \cdot 2}} & -\frac{1}{\sqrt{1 \cdot 2}} & 0 & \dots & 0 \\ \frac{1}{\sqrt{2 \cdot 3}} & \frac{1}{\sqrt{2 \cdot 3}} & -\frac{2}{\sqrt{2 \cdot 3}} & \ddots & \vdots \\ \vdots & & & \ddots & 0 \\ \frac{1}{\sqrt{(n-1) \cdot n}} & \dots & \frac{1}{\sqrt{(n-1) \cdot n}} & \frac{-(n-1)}{\sqrt{(n-1) \cdot n}} \\ \frac{1}{\sqrt{n}} & \dots & \dots & \frac{1}{\sqrt{n}} \end{pmatrix} = (\vec{c}_1, \dots, \vec{c}_n) = \begin{pmatrix} \vec{z}_1^T \\ \vdots \\ \vec{z}_n^T \end{pmatrix}$$

Z.B. • P.ex. $\vec{z}_j^T = (\underbrace{\frac{1}{\sqrt{(j) \cdot (j+1)}}, \frac{1}{\sqrt{(j) \cdot (j+1)}, \dots, \frac{1}{\sqrt{(j) \cdot (j+1)}}}_j, -\frac{j}{\sqrt{(j) \cdot (j+1)}}, 0, \dots, 0)$

Sei • Soit $j < k < n$

$$\begin{aligned} \leadsto \vec{z}_j^T \cdot \vec{z}_k &= \langle \vec{z}_j, \vec{z}_k \rangle = \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{(k) \cdot (k+1)}} + \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{(k) \cdot (k+1)}} + \dots \\ &+ \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{(k) \cdot (k+1)}} - \frac{j}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{(k) \cdot (k+1)}} + 0 = \frac{j-j}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{(k) \cdot (k+1)}} = 0 \end{aligned}$$

Sei • Soit $j = k < n \rightsquigarrow$

$$\begin{aligned} \vec{z}_j^T \cdot \vec{z}_j &= \left(\frac{1}{\sqrt{(j) \cdot (j+1)}}\right)^2 + \left(\frac{1}{\sqrt{(j) \cdot (j+1)}}\right)^2 + \dots + \left(\frac{1}{\sqrt{(j) \cdot (j+1)}}\right)^2 + \left(-\frac{j}{\sqrt{(j) \cdot (j+1)}}\right)^2 = \\ j \cdot \frac{1}{(j) \cdot (j+1)} + \frac{j^2}{(j) \cdot (j+1)} &= \frac{j^2 + 1}{j^2 + 1} = 1 \end{aligned}$$

Sei • Soit $j < k = n \rightsquigarrow \vec{z}_j^T \cdot \vec{z}_n = \langle \vec{z}_j, \vec{z}_n \rangle = \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{n}} + \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{n}} + \dots$
 $+ \frac{1}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{n}} - \frac{j}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{n}} = \frac{j-j}{\sqrt{(j) \cdot (j+1)}} \cdot \frac{1}{\sqrt{n}} = 0$

Sei • Soit $j = k = n \rightsquigarrow$

$$\vec{z}_n^T \cdot \vec{z}_n = \left(\frac{1}{\sqrt{n}}\right)^2 + \left(\frac{1}{\sqrt{n}}\right)^2 + \dots + \left(\frac{1}{\sqrt{n}}\right)^2 + n \cdot \frac{1}{n} = 1 \rightsquigarrow$$

$$\vec{z}_j^T \cdot \vec{z}_k = \delta_{i,k}, \quad M \cdot M^T = E \Rightarrow M^T = M^{-1} \Rightarrow \frac{1}{\det M^{-1}} = \det M = \det M^T = \det M^{-1} \Rightarrow \det M = \pm 1$$

Konsequenz: • **Conséquence:** $\rightsquigarrow M$ orthogonal: • *M orthogonale:*

$$\begin{aligned} \vec{Y} = M \cdot \vec{X} \Rightarrow r^2 = \vec{Y}^2 = \langle \vec{Y}, \vec{Y} \rangle &= \langle M \cdot \vec{X}, M \cdot \vec{X} \rangle = (M \cdot \vec{X})^T \cdot M \cdot \vec{X} = (\vec{X}^T \cdot M^T) \cdot (M \cdot \vec{X}) = \\ \vec{X}^T \cdot (M^T \cdot M) \cdot \vec{X} &= \vec{X}^T \cdot E \cdot \vec{X} = \vec{X}^T \cdot \vec{X} = \langle \vec{X}, \vec{X} \rangle = \vec{X}^2 \Rightarrow |\vec{X}| = |\vec{Y}|, \quad \sum_{i=1}^n X_i^2 = \sum_{i=1}^n Y_i^2 \end{aligned}$$

Nach Konstruktion war: • *D'après la construction on avait:*

$$\begin{aligned} Y_n &= \frac{1}{\sqrt{n}} (X_1 + X_2 + \dots + X_{n-1} + X_n) = \sqrt{n} \bar{X} \Rightarrow Y_n^2 = n \bar{X}^2 \rightsquigarrow \\ Y_n^2 &= \frac{1}{n} \left(\sum_{i=1}^n X_i \right)^2 = n \cdot \left(\frac{1}{n} \sum_{i=1}^n X_i \right)^2 = n \cdot (\bar{X})^2 \Rightarrow Q_n = \sum_{i=1}^n (X_i - \bar{X})^2 = \sum_{i=1}^n (X_i^2) - 2 X_i \bar{X} + \bar{X}^2 \\ &= \left(\sum_{i=1}^n (X_i^2) \right) - (2 \bar{X} \sum_{i=1}^n X_i) + \left(\sum_{i=1}^n \bar{X}^2 \right) = \left(\sum_{i=1}^n (X_i^2) \right) - (2 \bar{X} n \bar{X}) + (n \bar{X}^2) = \left(\sum_{i=1}^n (X_i^2) \right) - (n \bar{X}^2) \\ &= \left(\sum_{i=1}^n (Y_i^2) \right) - (Y_n^2) = \sum_{i=1}^{n-1} (Y_i^2) \end{aligned}$$

Dass die Y_i unabhängig sind, ersieht man aus der Definition der Y_i . Es ist $Y_{i-1} = Y_{i-1}(X_1, X_2, \dots, \cdot_i)$, $Y_i = Y_i(X_1, X_2, \dots, \cdot_{i+1})$. Daher ist es möglich, dass die eine Variable ändert, die andere aber nicht.

• *Les Y_i sont indépendantes comme on voit par la définition des Y_i . Il est $Y_{i-1} = Y_{i-1}(X_1, X_2, \dots, \cdot_i)$, $Y_i = Y_i(X_1, X_2, \dots, \cdot_{i+1})$. Il est donc possible de changer la valeur d'une variable sans toucher les valeurs des autres variables.* $\rightsquigarrow \text{☺}$

Lemma: • **Lemme:**

$$\sum_{i=1}^n X_i^2 = \sum_{i=1}^n Y_i^2 = \sum_{i=1}^{n-1} Y_i^2 + Y_n^2 = Q_n + n \bar{X}^2 = \sum_{i=1}^n (X_i - \bar{X})^2 + n \bar{X}^2 = \sum_{i=1}^{n-1} Y_i^2 + n \bar{X}^2, \quad Y_n^2 = n \bar{X}^2$$

Weiter verwenden wir den folgenden Satz von Seite 141: • *On outre nous utilisons le théorème suivant de la page 141:*

Satz: • **Théorème:** **Summe unabhängiger normalverteilter Variablen**
 • **Somme de variables indépendantes et distribuées de façon normale**

Vor.: • **Hyp.:**

Seien $X_1, X_2, \dots, X_n$ unabhängig, normalverteilt mit

• *Soient $X_1, X_2, \dots, X_n$ indépendantes, distribuées de façon normale avec*

Mittelwerte • *Moyennes $\mu_1, \mu_2, \dots, \mu_n$*

Varianzen • *Variances $\sigma_1^2, \sigma_2^2, \dots, \sigma_n^2$*

Beh.: • **Thè.:**

$\sum_{i=1}^n X_i$ normalverteilt • *distribuées de façon normale*

Mittelwert • *Moyenne $\mu = \sum_{i=1}^n \mu_i$*

Varianz • *Variance $\sigma^2 = \sum_{i=1}^n \sigma_i^2$*

Seien die X_i unabhängig und $\in N(\mu, \sigma^2)$, $k < n$. Dann sind auch die Summanden von Y_k , d.h. $X_1, X_2, \dots, -k X_{k+1}$ unabhängig und es gilt (vgl. Seite 158):

• *Les X_i soient indépendantes et $\in N(\mu, \sigma^2)$, $k < n$. Par conséquent les termes de la somme dans Y_k , ç.v.d. $\frac{X_1}{\sqrt{k \cdot (k+1)}}, \frac{X_2}{\sqrt{k \cdot (k+1)}}, \dots, -\frac{k X_{k+1}}{\sqrt{k \cdot (k+1)}}$ sont aussi indépendantes et il vaut (voir page 158):*

$$\mu(Y_k) = \sum_{i=1}^n \mu_i = \frac{\mu}{\sqrt{k \cdot (k+1)}} + \frac{\mu}{\sqrt{k \cdot (k+1)}} + \dots - \frac{k \mu}{\sqrt{k \cdot (k+1)}} = \mu \left(k \frac{1}{\sqrt{k \cdot (k+1)}} - \frac{k}{\sqrt{k \cdot (k+1)}} \right) = 0$$

$$\sigma^2(Y_k) = \sum_{i=1}^n \sigma_i^2 = \frac{\sigma^2}{k \cdot (k+1)} + \frac{\sigma^2}{k \cdot (k+1)} + \dots + (-k)^2 \frac{\sigma^2}{k \cdot (k+1)} = \sigma^2 \frac{k + k^2}{k^2 + k} = \sigma^2$$

$$\Rightarrow \sigma^2\left(\frac{Y_k}{\sigma}\right) = \left(\frac{1}{\sigma}\right)^2 \cdot \sigma^2(Y_k) = \frac{\sigma^2}{\sigma^2} = 1 \rightsquigarrow$$

Lemma: • **Lemme:** **Vor.:** • **Hyp.:**

$X_i \in N(\mu, \sigma^2)$, $X_i \in \{\heartsuit\}$, $k < n$,

$$Y_k = \frac{1}{\sqrt{k \cdot (k+1)}} (X_1 + X_2 + \dots + X_k - k X_{k+1})$$

Beh.: • **Thè.:**

$$\frac{Y_k}{\sigma} \in N(0, 1)$$

$$\text{Sei} \quad \bullet \quad \text{Soit} \quad Y := \frac{Q_n}{\sigma^2} = \frac{1}{\sigma^2} \sum_{i=2}^{n-1} Y_i^2 = \sum_{i=2}^{n-1} \frac{Y_i^2}{\sigma^2} = \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 \rightsquigarrow$$

Y ist eine Summe von $n - 1$ Quadraten von normalverteilten Variablen $\frac{Y_i}{\sigma} \in N(0, 1)$. Die Y_i sind unabhängig. Y ist somit χ^2 -verteilt mit $n - 1$ Freiheitsgraden! Man hat somit den Satz:

• *Y est une Somme de $n - 1$ variables au carré $\frac{Y_i}{\sigma} \in N(0, 1)$ qui sont distribuées de façon normale. Les Y_i sont indépendantes. Y est donc réparti selon la loi χ^2 avec $n - 1$ degrés de liberté! On a donc le théorème:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\begin{aligned} X_1, \dots, X_n &\in N(\mu, \sigma^2), \in \{\heartsuit\} \\ \bar{X} &= \frac{1}{n} \cdot \sum_{i=1}^n X_i \\ Y &= \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 \end{aligned}$$

Beh.: • **Thè.:**

$$Y \in \chi_{n-1}^2$$

Sei • *Soit* $Z = \sqrt{n} \frac{\bar{X} - \mu}{\sigma} \rightsquigarrow$ Es gilt: • *Il vaut:*

$$Z = \frac{\sqrt{n}}{\sigma} \cdot \bar{X} + \frac{(-1) \cdot \sqrt{n} \cdot \mu}{\sigma} = c_1 \bar{X} + c_2 \Rightarrow (Z \in \{\heartsuit\} \Leftrightarrow \bar{X} \in \{\heartsuit\})$$

Weiter sieht man: • *En plus on a:*

$$\begin{aligned} Q_n &= \sum_{i=1}^n (X_i - \bar{X})^2 \Rightarrow \frac{\partial}{\partial \bar{X}} Q_n = \frac{\partial}{\partial \bar{X}} \sum_{i=1}^n (X_i - \bar{X})^2 = 2 \sum_{i=1}^n (X_i - \bar{X}) \cdot (-1) = -2 \left(\left(\sum_{i=1}^n X_i \right) - n \cdot \bar{X} \right) \\ &= -2 \left((n \cdot \bar{X}) - n \cdot \bar{X} \right) = 0 \Rightarrow Q_n(\bar{X}) = \text{const.}(X_1, \dots, X_n) \end{aligned}$$

Andererseits ist: • *D'autre part on voit:*

$k < n - 1 \rightsquigarrow Y_k = Y_k(X_1, \dots, X_{k+1}), \bar{X} = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}(X_1, \dots, X_n) \rightsquigarrow$ Man kann Y_k verändern ohne $\bar{X}$ zu verändern und umgekehrt. • *On peut changer le Y_k sans changer le $\bar{X}$ et vice versa.*

$k = n - 1 \rightsquigarrow Y_{n-1} = \lambda_1 \cdot \bar{X}(X_1, \dots, X_n) + \lambda_2 \cdot X_n \rightsquigarrow$ Auch hier kann man Y_k verändern ohne $\bar{X}$ zu verändern und umgekehrt. • *On peut aussi changer le Y_k sans changer le $\bar{X}$ et vice versa.* $\rightsquigarrow$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\begin{aligned} X_1, \dots, X_n &\in N(\mu, \sigma^2), \in \{\heartsuit\} \\ \bar{X} &= \frac{1}{n} \cdot \sum_{i=1}^n X_i \\ Y &= \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 \\ Z &= \sqrt{n} \frac{\bar{X} - \mu}{\sigma} \end{aligned}$$

Beh.: • **Thè.:**

$$\{Y, Z\} \subset \{\heartsuit\}$$

Konsequenz: • Conséquence:

Mit $Z \in N(0, 1)$ und $Y \in \chi_{n-1}^2$ lässt sich nun die Zufallsvariable $T = \frac{Z}{\sqrt{Y/(n-1)}}$ bilden, die somit einer t -Verteilung resp. einer Student-Verteilung genügt.

• Avec $Z \in N(0, 1)$ et $Y \in \chi_{n-1}^2$ on peut construire la variable aléatoire $T = \frac{Z}{\sqrt{Y/(n-1)}}$ qui suit donc la loi de la répartition de Student.

Daher folgt: • On déduit donc:

$$T = \sqrt{n} \frac{\bar{X} - \mu}{\sigma} \cdot \frac{1}{\sqrt{\frac{1}{\sigma^2} \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}} = \sqrt{n} \frac{\bar{X} - \mu}{1} \cdot \frac{1}{\sqrt{S^2}} = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$$

Korollar: • Corollaire:**Vor.: • Hyp.:**

$$X_1, \dots, X_n \in N(\mu, \sigma^2), \in \{\heartsuit\}$$

$$\bar{X} = \frac{1}{n} \cdot \sum_{i=1}^n X_i$$

$$S = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$$

$$T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$$

Beh.: • Thè.:

$$T \in t_{n-1}$$

Bemerkung: • Remarque:

Man beachte, dass wir von $T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$ jetzt die Verteilungsfunktion kennen. In T steckt neben den Stichprobenwerten X_i aber nur μ . Man hat daher hier ein Mittel in der Hand anhand von Stichprobenwerten zu einer Aussage über μ zu gelangen! • *Considérons que nous connaissons maintenant la fonction de répartition de $T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$. Dans T il y a à côté des valeurs d'échantillon X_i seulement le μ . On a donc ici un moyen en main qui peut servir à obtenir de l'information sur le μ à partir des valeurs de l'échantillon!*

5.5.6 Vertrauensintervalle II — Intervalles de confiance II

Problem: • Problème: Wegen der Konsistenz wissen wir, dass bei grossen Stichproben $\bar{X}$ ein guter Schätzer für μ ist. Wie ist es aber bei kleinen Stichproben?

• À cause de la consistance, nous savons que pour les grandes échantillons, $\bar{X}$ est une bonne estimation de μ . Mais comment est-ce que se présente la chose pour les échantillons petits?

Mittelwertschätzung bei unbekannter Varianz und kleinen Stichproben — Estimation de la moyenne à la variance inconnue et des petits échantillons

Gegeben seien die Stichprobenwerte $\bar{X} = \frac{1}{n} \cdot \sum_{i=1}^n X_i$, $S = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$ und n . Damit bilden

wir die Variable $T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$, die nun einer t -Verteilung mit $n - 1$ Freiheitsgraden genügt. Wie kennen die zugehörigen Funktionen:

• Soient données les valeurs d'un échantillon $\bar{X} = \frac{1}{n} \cdot \sum_{i=1}^n X_i$, $S = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}$ et n . Nous

en construisons la variable $T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$ qui suit la loi de Student à $n - 1$ degrés de liberté. Nous connaissons les fonctions liées:

$$f(z) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n} \pi \Gamma(\frac{n}{2})} \cdot \frac{1}{(1 + \frac{z^2}{n})^{(n+1)/2}}, \quad F(z) = P(T \leq z) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n} \pi \Gamma(\frac{n}{2})} \cdot \int_{-\infty}^z \frac{1}{(1 + \frac{u^2}{n})^{(n+1)/2}} du$$

Aus $f(z)$ ersieht man, dass die Verteilung symmetrisch ist zur y -Achse. Daher gilt:

• On voit de $f(z)$ que la fonction de distribution est symétrique à l'axe y . Il vaut donc:

$$F(-c) = 1 - F(c) \Rightarrow P(-c \leq T \leq c) = P(T \leq c) - P(T < -c) = F(c) - F(-c) = F(c) - (1 - F(c)) = 2F(c) - 1$$

Für eine vorgegebene Wahrscheinlichkeit $P(-c \leq T \leq c) = \gamma = 1 - \alpha$ kann man daher aus $\gamma = 2F(c) - 1$ das c bestimmen.

• Pour une probabilité donnée $P(-c \leq T \leq c) = \gamma = 1 - \alpha$ on peut donc calculer le c de l'équation $\gamma = 2F(c) - 1$.

Mit Hilfe des nun bekannten c kann man jetzt auf ein Vertrauensintervall für μ schließen: • A l'aide d'une valeur c connue on peut maintenant en conclure l'extension de l'intervalle de confiance:

$$-c \leq T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S} \leq c \Rightarrow \frac{-c \cdot S}{\sqrt{n}} - \bar{X} \leq -\mu \leq \frac{c \cdot S}{\sqrt{n}} - \bar{X} \Rightarrow \bar{X} - \frac{c \cdot S}{\sqrt{n}} \leq \mu \leq \bar{X} + \frac{c \cdot S}{\sqrt{n}},$$

$$I = [\bar{X} - \frac{c \cdot S}{\sqrt{n}}, \bar{X} + \frac{c \cdot S}{\sqrt{n}}]$$

Methode: • **Méthode:**

1. **Geg.:** • **Donné:** $\{x_1, x_2, \dots, x_n\}$, $\gamma = 1 - \alpha$

$$2. \rightsquigarrow \bar{X} = \frac{1}{n} \cdot \sum_{i=1}^n X_i, \quad S = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2}, \quad T = \sqrt{n} \cdot \frac{\bar{X} - \mu}{S}$$

$$3. \rightsquigarrow \gamma = 2F(c) - 1 \rightsquigarrow c$$

$$4. \rightsquigarrow \mu \in I = [\bar{X} - \frac{c \cdot S}{\sqrt{n}}, \bar{X} + \frac{c \cdot S}{\sqrt{n}}] \text{ mit Wahrscheinlichkeit } \gamma \quad \bullet \text{ avec la probabilité } \gamma$$

Zahlenbeispiel — Exemple avec des nombres

Sei • *Soit* $\gamma = 0.95$, $\alpha = 0.05$ Daten erzeugen: • *Créer les données:*

Mathematica-Programm: • Programme en Mathematica:

```
<< Graphics'Graphics';
<< Statistics'DescriptiveStatistics';
gamma = 0.95;
m = Table[243.79 + 2.9 Random[], {n, 1, 25}]
```

Output: • Output:

```
{246.21, 244.807, 244.481, 245.184, 246.195, 245.083, 244.915, 246.492,
243.98, 243.993, 246.646, 244.568,
245.833, 243.864, 246.659, 244.947,
246.081, 243.834, 245.631, 244.54, 246.166, 245.966, 245.905, 246.605,
246.645}
```

Mathematica-Programm: • Programme en Mathematica:

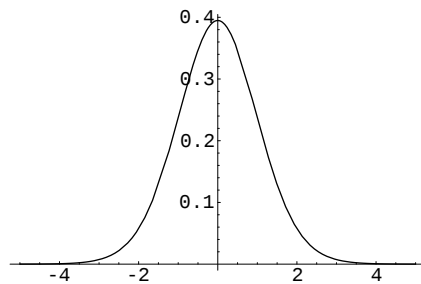
```
meanM = Mean /. LocationReport[m][[1]];
standD = StandardDeviation /. DispersionReport[m][[2]];
{LocationReport[m], DispersionReport[m], Length[m]}
```

Output: • Output:

```
{{Mean -> 245.409, HarmonicMean -> 245.406, Median -> 245.631},
{Variance -> 0.93796, StandardDeviation -> 0.968484, SampleRange -> 2.82473,
MeanDeviation -> 0.8572, MedianDeviation -> 0.823889,
QuartileDeviation -> 0.819016}, 25}
```

Mathematica-Programm: • Programme en Mathematica:

```
f[z_, m_] :=
  Gamma[(m + 1)/2]/(Sqrt[m Pi] Gamma[m/2]) / (1 + z^2/m)^((m + 1)/2);
f[u, Length[m] - 1];
Plot[f[u, Length[m] - 1], {u, -5, 5};
```



Der Plot zeigt die Dichte der t -Verteilung $f(z, 24)$.

- *On voit dans l'esquisse la densité de la loi de Student $f(z, 24)$.*

Nun berechnen wir c direkt mit Mathematica. Wer für eine solche Berechnung die Mittel nicht hat, muss auf **Tabellen** zurückgreifen!

- *Maintenant nous calculons c directement avec Mathematica. Ceux qui n'ont pas les moyens pour un tel calcul, doivent avoir recours tableaux!*

Mathematica-Programm: • Programme en Mathematica:

```
{ "c", c = c /. FindRoot[gamma == 2 Evaluate[Integrate[f[u, 15],
{u, -Infinity, c}]] - 1, {c, 2}] // Chop,
"Interv", {xU = meanM - standD c/Sqrt[Length[m]],
x0 = meanM + standD c/Sqrt[Length[m]]} }
```

Output: • Output:

```
{ "c", 2.13145, "Interv", {244.996, 245.822} }
```

Mathematica-Programm: • Programme en Mathematica:

```
m1 = Select[m, (xU - 0.000001 < #1 < x0 + 0.000001) &]
```

Output: • Output:

```
{245.184, 245.083, 245.631}
```

Mathematica-Programm: • Programme en Mathematica:

```
Complement[m, m1]
```

Output: • Output:

{243.834, 243.864, 243.98, 243.993, 244.481, 244.54, 244.568, 244.807, 244.915, 244.947, 245.833, 245.905, 245.966, 246.081, 246.166, 246.195, 246.21, 246.492, 246.605, 246.645, 246.646, 246.659}

Lösung: • **Solution:** Wie man sieht, liegt der Mittelwert μ mit einer Wahrscheinlichkeit von 95% im Intervall $I = [244.996, 245.822]$ ($\bar{X} \approx 245.409$). Die meisten der Daten liegen jedoch ausserhalb dieses Intervalls. Das zeigt, wie gut der Schätzer $\bar{X}$ für μ ist.

• *Comme on peut remarquer, la moyenne μ se trouve avec une probabilité de 95% dans l'intervalle $I = [244.996, 245.822]$ ($\bar{X} \approx 245.409$). Par contre la plupart des données se trouve à l'extérieur de cet intervalle. Ça nous montre la qualité de l'estimation $\bar{X}$ de μ .*

5.5.7 Vertrauensintervalle III — Intervalles de confiance III

Problem: • **Problème:** Wegen der Konsistenz wissen wir, dass wir bei grossen Stichproben S^2 ein guter Schätzer für σ^2 ist. Wie ist es aber bei kleinen Stichproben?

• *À cause de la consistance, nous savons que pour les grandes échantillons, S^2 est une bonne estimation de σ^2 . Mais comment est-ce que se présente la chose pour les petits échantillons?*

Varianzschätzung bei unbekanntem Mittelwert und kleinen Stichproben — Estimation de la variance à la moyenne inconnue et les petits échantillons

Auf Seite 188 haben wir festgestellt: • *A la page 188 on a constaté:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

$$\begin{aligned} X_1, \dots, X_n &\in N(\mu, \sigma^2), \in \{\heartsuit\} \\ \bar{X} &= \frac{1}{n} \cdot \sum_{i=1}^n X_i \\ Y &= \frac{1}{\sigma^2} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{n-1}{\sigma^2} S^2 \end{aligned}$$

Beh.: • **Thè.:**

$$Y \in \chi_{n-1}^2$$

Konsequenz: • **Conséquence:**

$$\gamma = P(c \leq Y) = 1 - P(0 \leq Y < c) = 1 - (F_{\chi^2, n-1}(c) - F_{\chi^2, n-1}(0)) = 1 - F_{\chi^2, n-1}(c)$$

Bemerkung: • **Remarque:**

Da die Varianz stets positiv ist, begnügen wir uns hier mit einer einseitigen Abschätzung im Gegensatz zu der zweiseitigen Abschätzung beim Vertrauensintervall für μ .

• *Comme la variance est toujours positive, nous nous contentons ici d'une estimation seulement dans une direction, par contre à l'estimation dans deux directions pour l'intervalle de confiance pour μ .*

Wir wissen: • *Nous savons:*

$$f(x) = \begin{cases} 0 & x \leq 0 \\ K_n x^{\frac{n-2}{2}} e^{-\frac{x}{c}} & x > 0 \end{cases} \quad K_n := \frac{1}{2^{\frac{n}{2}} \Gamma(\frac{n}{2})}$$

$$F(x) = K_n \int_0^x u^{\frac{n-2}{2}} e^{-\frac{u}{c}} du$$

Damit können wir c berechnen aus $\gamma = P(c \leq Y) = 1 - F_{\chi^2, n-1}(c)$.

• *Nous pouvons donc calculer c de $\gamma = P(c \leq Y) = 1 - F_{\chi^2, n-1}(c)$.*

$$\leadsto c \leq Y = \frac{n-1}{\sigma^2} S^2 \Rightarrow 0 \leq \sigma^2 \leq \frac{n-1}{c} S^2, \sigma^2 \in [0, \frac{n-1}{c} S^2]$$

Methode: • **Méthode:**

1. **Geg.:** • **Donné:** $\{x_1, x_2, \dots, x_n\}$, $\gamma = 1 - \alpha$

$$2. \leadsto S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2, Y = \frac{n-1}{\sigma^2} S^2$$

$$3. \leadsto \gamma = 1 - F_{\chi^2, n-1}(c) \leadsto c$$

$$4. \leadsto \sigma^2 \in [0, \frac{n-1}{c} S^2] \text{ mit Wahrscheinlichkeit } \gamma \quad \bullet \text{ avec la probabilité } \gamma$$

Zahlenbeispiel — Exemple avec des nombres

Sei • *Soit* $\gamma = 0.95$, $\alpha = 0.05$ Daten erzeugen: • *Créer les données:*

Mathematica-Programm: • **Programme en Mathematica:**

```
<< Graphics'Graphics';
<< Statistics'DescriptiveStatistics';
gamma = 0.95;
m = Table[243.79 + 2.9 Random[], {n, 1, 25}]
```

Output: • **Output:**

```
{244.584, 244.288, 245.634, 245.936, 244.285, 245.569, 245.374, 246.273,
243.942, 244.729, 244.258, 246.685, 246.468, 245.045, 245.432, 244.24,
245.035, 246.111, 246.083, 243.971, 245.717, 244.192, 246.688, 244.861, 244.923}
```

Mathematica-Programm: • **Programme en Mathematica:**

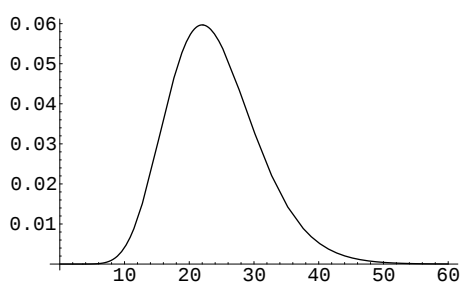
```
var = Variance /. DispersionReport[m][[1]];
{LocationReport[m], DispersionReport[m], Length[m]}
```

Output: • **Output:**

```
{Mean -> 245.213, HarmonicMean -> 245.21, Median -> 245.045},
{Variance -> 0.768021, StandardDeviation -> 0.876368, SampleRange -> 2.74647,
MeanDeviation -> 0.753255, MedianDeviation -> 0.760244,
QuartileDeviation -> 0.842562}, 25}
```

Mathematica-Programm: • Programme en Mathematica:

```
chi[x_, n_] := 1/(2^(n/2) Gamma[n/2]) x^((n - 2)/2) E^(-x/2);
Plot[chi[u, Length[m] - 1], {u, 0, 60}];
```



Der Plot zeigt die Dichte der χ^2 -Verteilung $chi(x, 24)$.

• On voit dans l'esquisse la densité de la loi de khi deux $chi(x, 24)$.

Nun berechnen wir c direkt mit Mathematica. Wer für eine solche Berechnung die Mittel nicht hat, muss auf **Tabellen** zurückgreifen!

• *Maintenant nous calculons c directement avec Mathematica. Ceux qui n'ont pas les moyens pour un tel calcul, doivent avoir recours aux tableaux!*

Mathematica-Programm: • Programme en Mathematica:

```
{"c", c = c /. FindRoot[gamma == 1 - Evaluate[Integrate[chi[u, 15],
{u, 0, c}]], {c, 40}] // Chop, "Interv", {xU = 0, x0 = ((Length[m]-1)/c) var } }
```

Output: • Output:

```
{"c", 7.26094, "Interv", {0, 2.53858}}
```

Lösung: • **Solution:** Wie man sieht, liegt die Varianz σ^2 mit einer Wahrscheinlichkeit von 95% im Intervall $I = [0, 2.53858]$ ($S^2 \approx 0.768021$). Weitere „Runs“ liefern: $\gamma = 0.99 \rightsquigarrow I = [0, 3.52486]$, $\gamma = 0.90 \rightsquigarrow I = [0, 2.15667]$

• *Comme on peut remarquer, la variance σ^2 se trouve avec une probabilité de 95% dans l'intervalle $I = [0, 2.53858]$ ($S^2 \approx 0.768021$). Par d'autres "runs" on obtient: $\gamma = 0.99 \rightsquigarrow I = [0, 3.52486]$, $\gamma = 0.90 \rightsquigarrow I = [0, 2.15667]$*

5.5.8 Vertrauensintervalle IV — Intervalles de confiance IV

Bemerkung zur Binomialverteilung — Remarque concernant la loi binomiale

Wahrscheinlichkeitsfunktion: • *Fonction de probabilité:*

$$f(x) = P(X = x) = \binom{n}{x} p^x q^{n-x}, \quad q = 1 - p, \quad \mu = np, \quad \sigma^2 = npq$$

Problem: • **Problème:** Gesucht ist ein Vertrauensintervall für p bei der Binomialverteilung.

• *On cherche un intervalle de confiance pour p de la répartition binomiale.*

Für grosse n gilt die Näherungsformel: • *Pour des n grands la formule d'approximation est valable:*

$$f(x) \approx f^*(x) = \frac{1}{\sqrt{w\pi}\sigma} \cdot e^{-\frac{z^2}{2}}, \quad Z = \frac{X - \mu}{\sigma}, \quad Z \in N(0, 1)$$

Seien X die Anzahl Erfolge beim n -maligen Ausführen des Experimentes. Dann kann man wieder wie folgt argumentieren: • *Soit X le nombre de succès à l'occasion de n répétitions de l'expérience. Puis on peut argumenter comme il suit:*

$$\gamma = P(-c \leq Z \leq c) \approx \int_{-c}^c f^*(x) dx$$

$$\Rightarrow c = \dots \Rightarrow -c \leq Z = \frac{X - \mu}{\sigma} \leq c \Rightarrow -c \cdot \sigma + \mu \leq X \leq +c \cdot \sigma + \mu \dots$$

5.5.9 Automatische Bestimmung der Vertrauensbereiche — Calcul automatique des domaines de confiance

Dem geeigneten Leser wird es wohl klar sein, dass heutzutage für alle Standardverfahren Computerprogramme vorhanden sind. Mit Mathematica (z.B. V.4.0) kann man Vertrauensbereiche wie folgt bestimmen:

• *Pour le lecteur moderne, il est certainement clair qu'aujourd'hui on a des programmes d'ordinateur à disposition pour toutes les méthodes standard. Avec Mathematica (par exemple V.4.0) on calcule des domaines de confiance comme il suit:*

Bsp.: • Exemple:

Vertrauensbereiche für Mittelwert und Varianz, Vertrauensniveau 0.9. Als Grundlage verwendete Verteilungen: t -Verteilung für den Mittelwert (mean) und χ^2 -Verteilung für die Varianz (var).

• *Intervalles de confiance pour la moyenne et la variance, niveau de confiance 0.9. Lois utilisées comme base: Loi de Student pour la moyenne (mean) et loi de χ^2 pour la variance (var).*

Mathematica-Programm: • Programme en Mathematica:

```
<< Statistics'ConfidenceIntervals';
data = {2.2, 1.3, 0.8, 1.0, 1.1, 3.1, 3.4, 3.3, 2.2, 0.5};
{"mean->", MeanCI[data, ConfidenceLevel -> 0.9],
"var->", VarianceCI[data, ConfidenceLevel -> 0.9]}
```

Output: • Output:

```
| {"mean->", {1.25473, 2.52527}, "var->", {0.638868, 3.25072}} |
```

5.6 Signifikanztests — Tests de signification

5.6.1 Hypothesen — Hypothèses

Idee: • **Idée:** Einen Parameter testen heisst, zuerst über die Lage dieses Parameters **Hypothesen** aufzustellen und dann anschliessend die Wahrscheinlichkeit dieser Hypothesen zu untersuchen. Üblicherweise führt das dann zu einer Verwerfung oder im andern Fall zu einer Tolerierung resp. einer Akzeptanz der Hypothesen auf der Grundlage einer geforderten Wahrscheinlichkeit. Wie das genau funktioniert, wollen wir an charakteristischen Beispielen studieren.

• *Tester un paramètre signifie de postuler d'abord des **hypothèses** concernant la position du paramètre en question et ensuite d'étudier la probabilité de l'hypothèse. Normalement ce processus mène à un refus ou dans l'autre cas à la tolérance resp. à accepter l'hypothèse sur la base d'une probabilité demandée. Nous voulons étudier au moyen d'exemples caractéristiques, comment ça fonctionne exactement.*

Methode: • **Méthode:**

- Wir beginnen mit der **Nullhypothese** H_0 : Z.B. postulieren wir, dass ein Parameter z einen Sollwert z_0 einhält: $H_0 = (z = z_0)$. Z.B. $H_0 = (\mu = \mu_0)$
 - *Nous commençons par l'hypothèse nulle H_0 : P.ex. nous postulons qu'un paramètre z respecte une valeur de consigne z_0 : $H_0 = (z = z_0)$. P.ex. $H_0 = (\mu = \mu_0)$*

Andere Beispiele: • *Autres exemples:*

$$H_0 = (z \neq z_0) \Leftrightarrow H_0 = (z \leq z_0), \quad H_0 = (z \not\leq z_0) \Leftrightarrow H_0 = (z \geq z_0)$$

- Normalerweise wird der **Nullhypothese**, für die eine **hohe Wahrscheinlichkeit** zu erwarten ist, eine **Alternativhypothese** H_1 mit **kleiner Wahrscheinlichkeit** α gegenübergestellt, z.B. hier $H_1 = (z \neq z_0)$. Jetzt kann man die Wahrscheinlichkeit der Abweichung $z - z_0$ mittels einer geeigneten Testverteilung feststellen. (Mit z_0 fliesst somit auch die Nullhypothese als Voraussetzung in das Verfahren ein.) Man verwendet dazu normalerweise einen geeigneten Schätzer $\zeta(z)$ für z bei einem gegebenen Stichprobenumfang. Falls der real vorhandene Wert bei nur sehr kleiner verrechneten Wahrscheinlichkeit α trotzdem für die Alternativhypothese spricht, so hat man es mit einem seltenen Ereignis zu tun. Die Zufallsabweichung ist daher unwahrscheinlich. Der Alternativhypothese muss daher eine grössere als die angenommene Wahrscheinlichkeit zugebilligt werden, die Nullhypothese ist daher zu verwerfen. Im andern Fall kann man die kleine Wahrscheinlichkeit der Alternativhypothese nicht angreifen und somit die Nullhypothese nicht verwerfen. Man muss diese also tolerieren resp. man akzeptiert sie („auf dem Niveau α “).
 - *Normalement on oppose une **hypothèse alternative** H_1 avec une **petite probabilité** α à l'hypothèse nulle avec une **grande probabilité**, p.ex. ici $H_1 = (z \neq z_0)$. Maintenant on peut étudier la probabilité de la différence $z - z_0$ par le moyen d'une répartition de test convenable. (Avec z_0 , l'hypothèse nulle entre ici comme condition aussi dans la procédure.) Normalement on utilise une estimation convenable $\zeta(z)$ pour z et on a une taille d'échantillon donnée. Si sous la condition d'une probabilité α petite qu'on a utilisée pour calculer, la valeur qui est actuellement présente favorise quand-même l'hypothèse alternative, on a à faire à un événement rare. Une différence aléatoire est donc improbable. A l'hypothèse alternative il faut donc attribuer une probabilité qui est plus grande que celle (α) qu'on vient d'utiliser. Par conséquent on repousse l'hypothèse nulle. Dans l'autre cas, on ne peut pas attaquer la petite probabilité de l'hypothèse alternative. Par conséquent*

on ne peut pas repousser l'hypothèse nulle. On doit la tolérer ainsi resp. on l'accepte ("au niveau α ").

5.6.2 Zweiseitige Alternative, t -Test — Alternative bilatérale, test de Student

Bsp.: • **Exemple:** **Geg.:** • **Donné:**

$X \in N(\mu, \sigma^2)$, $H_0 = (\mu = \mu_0) \rightsquigarrow H_1 = (\mu \neq \mu_0)$, $\zeta(\mu) = \bar{x} = \bar{x}_n$

μ_0 ist ein abgemachter, durch Pläne vorgegebener und einzuhaltender Wert.

• μ_0 est une valeur qu'il faut respecter selon l'accord et qui est donnée par les plans.

Sei • Soit $\bar{x} = 127.3$, $\mu_0 = 128$, $s = 1.2$, $n = 20$, $\alpha = 0.01$

Dabei kann man als **Testverteilung** die Stichprobenfunktion „normierte Differenz“ oder Abweichung von μ_0 verwenden. Hier wird als Voraussetzung die Nullhypothese als in Form des Terms $\bar{X} - \mu_0$ eingebaut. μ_0 wird bei der Normierung auf null abgebildet.

• On peut utiliser comme **distribution de test** la fonction de l'échantillon "différence normée" ou l'écart de μ_0 . Ici on insère comme base aussi l'hypothèse nulle en forme du terme $\bar{X} - \mu_0$. μ_0 est appliqué à zéro lors de la normalisation.

$$\rightsquigarrow T = \frac{\bar{X} - \mu_0}{S/\sqrt{n}} = \sqrt{n} \frac{\bar{X} - \mu_0}{S}$$

Wichtig: • **Important:** $\rightsquigarrow T$ genügt einer t -Verteilung mit $n - 1$ Freiheitsgraden.

• $\rightsquigarrow T$ satisfait une répartition t avec $n - 1$ degrés de liberté.

$$T \rightsquigarrow f(T) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n} \pi \Gamma(\frac{n}{2})} \cdot \frac{1}{(1 + \frac{T^2}{n})^{\frac{n+1}{2}}}, \quad F(c) = \frac{\Gamma(\frac{n+1}{2})}{\sqrt{n} \pi \Gamma(\frac{n}{2})} \cdot \int_{-\infty}^c \frac{1}{(1 + \frac{u^2}{n})^{\frac{n+1}{2}}} du$$

$\rightsquigarrow$ Studiere: • *Etudier:* $P(|T| > c) = \alpha$

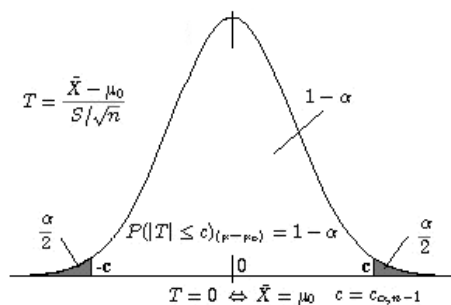
Die Wahrscheinlichkeit, dass unter der Voraussetzung $\mu = \mu_0$ die normierte Differenz grösser ist als c soll hier gleich einer gegebenen Wahrscheinlichkeit α sein. α ist eine Zahl, die man vernünftig festsetzen muss.

• La probabilité que sous la condition $\mu = \mu_0$ la différence normalisée soit plus grande que c , est ici égale à une probabilité donnée α . α est donc un nombre qu'il faut fixer de façon raisonnable.

Definition: • **Définition:**

α heisst **Signifikanzniveau** (Bedeutungsniveau).

• α s'appelle **niveau de signification**.



Es gilt: • Il vaut:

$$P(|T| > c) = \alpha \Leftrightarrow P(|T| \leq c) = 1 - \alpha,$$

$$P(|T| \leq c) = P(-c \leq T \leq c)$$

$$= F(c) - F(-c) = F(c) - (1 - F(c)) = 2F(c) - 1,$$

$$1 - \alpha = 2F(c) - 1 \Rightarrow 2 - 2F(c) = \alpha$$

$$1 - F(c) = \frac{\alpha}{2} \rightsquigarrow c = \dots \rightsquigarrow$$

$$P(|T| > c) = \alpha \Rightarrow$$

$$(|T_{real}| > c)_{P=\alpha} \Leftrightarrow (\neg(|T_{real}| \leq c))_{P=\alpha}$$

$$\Leftrightarrow \neg((-c \leq T_{\text{real}} = \sqrt{n} \frac{\bar{X} - \mu_0}{S} \leq c))_{P=\alpha} \Leftrightarrow (\neg(-\frac{c \cdot s}{\sqrt{n}} + \mu_0 \leq \bar{X} \leq \frac{c \cdot s}{\sqrt{n}} + \mu_0))_{P=\alpha}$$

$$\Leftrightarrow ((\bar{X} < -\frac{c \cdot s}{\sqrt{n}} + \mu_0) \vee (\frac{c \cdot s}{\sqrt{n}} + \mu_0 < \bar{X}))_{P=\alpha}$$

Nun berechnen wir c direkt mit Mathematica. Wer für eine solche Berechnung die Mittel nicht hat, muss auf **Tabellen** zurückgreifen!

• *Maintenant nous calculons c directement avec Mathematica. Ceux qui n'ont pas les moyens pour un tel calcul, doivent avoir recours à des tableaux!*

Mathematica-Programm: • **Programme en Mathematica:**

```
alpha = 0.01;
f[z_, n_] :=
  Gamma[(n + 1)/2]/(Sqrt[n Pi] Gamma[n/2]) / (1 + z^2/n)^(n + 1)/2);
{"c", c = c /. FindRoot[
  alpha/2 == Evaluate[1 - Integrate[f[u, 20], {u, -Infinity, c}]], {c, 2}] //
Chop, "Interv", {-c, +c}]
```

Output: • **Output:**

```
{"c", 2.84534, "Interv", {-2.84534, 2.84534}}
```

$$((\bar{X} < -\frac{c \cdot s}{\sqrt{n}} + \mu_0) \vee (\frac{c \cdot s}{\sqrt{n}} + \mu_0 < \bar{X}))_{P=\alpha}$$

$$\leadsto ((127.3 < -\frac{2.84534 \cdot 1.2}{\sqrt{20}} + 128) \vee (\frac{2.84534 \cdot 1.2}{\sqrt{20}} + 128 < 127.3))_{P=0.01}$$

$$\Leftrightarrow ((127.3 < 127.237) \vee (128.763 < 127.3))_{P=0.01} \Leftrightarrow (127.3 \notin [127.237, 128.763])_{P=0.01}$$

Ergebnis: Mit einer Wahrscheinlichkeit von $\alpha = 0.01$ soll also 127.3 nicht im Intervall $[127.237, 128.763]$ liegen. 127.3 liegt aber in diesem Intervall. Das bedeutet, dass das so ist, weil wirklich die Nullhypothese $H_0 = (\mu = \mu_0)$ richtig ist oder weil das zufällig so ist. **Man kann daher hier die Nullhypothese auf Grund des Tests bei einem Signifikanzniveau $\alpha = 0.01$ nicht ablehnen.** Würde jedoch 127.3 ausserhalb des Intervalls liegen, so wäre die Wahrscheinlichkeit dafür 0.01. Bei unserer Testphilosophie müsste man die Nullhypothese also ablehnen.

• *Résultat: Ainsi avec une probabilité de $\alpha = 0.01$, 127.3 ne devrait pas être trouvé dans l'intervalle $[127.237, 128.763]$. Mais 127.3 se trouve justement dans cet intervalle. Ça signifie que c'est comme ça parce que l'hypothèse nulle $H_0 = (\mu = \mu_0)$ est correcte ou parce que c'est comme ça par hasard. Ici on ne peut donc pas refuser l'hypothèse nulle sur la base de ce test et le niveau de signification $\alpha = 0.01$. Cependant si 127.3 était à l'extérieur de l'intervalle, la probabilité serait 0.01. Sur la base de notre philosophie de test il faudrait donc refuser l'hypothèse nulle.*

Achtung: • **Attention:**

Kleine Wahrscheinlichkeit bedeutet nicht Unmöglichkeit!

• *Petite probabilité ne signifie pas impossibilité*

Bemerkung: • **Remarque:**

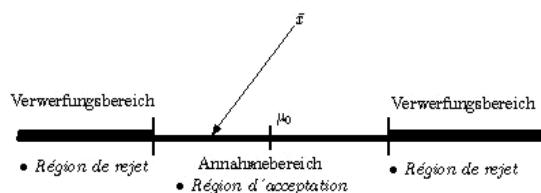
Das hier verwendete Testverfahren nennen wir **t-Test**, da wir die t -Verteilung von Student als Testverteilung verwendet haben.

• *La méthode de test que nous venons d'utiliser ici s'appelle **test de Student** parce que nous avons utilisé la loi de Student comme fonction de distribution.*

Annahmebereich, Verwerfungsbereich — région d'acceptation, domaine de rejet

Im eben besprochenen Beispiel ist die Alternativhypothese verworfen worden, weil der berechnete Wert des Schätzers im gefundenen Intervall und nicht ausserhalb lag. Wir reden hier von **Annahmebereich** und **Verwerfungsbereichen** (vgl. Skizze).

• Dans l'exemple qu'on vient de discuter, l'hypothèse alternative a été repoussée parce que la valeur calculée de l'estimation se trouvait dans l'intervalle trouvé et pas à l'extérieur de l'intervalle. Nous parlons ici du **domaine d'acceptation** et des **domaines de refus** (voir esquisse).



Achtung: • Attention:

Die Verwerfungsbereiche müssen nicht in allen Fällen symmetrisch liegen. Sie können auch einseitig liegen!

• Ce n'est pas la règle que les domaines de rejet soient symétriques. Ils peuvent se trouver aussi seulement d'un côté.

5.6.3 Einseitige Alternative, t -Test — Alternative unilatérale, test de Student

Wir betrachten hier den Fall $H_0 = (z = z_0)$ contra $H_1 = (z > z_0) \rightsquigarrow$ **einseitiger Alternativtest**.

• Nous considérons ici le cas $H_0 = (z = z_0)$ contra $H_1 = (z > z_0) \rightsquigarrow$ **test alternatif unilatéral**.

Bsp.: • Exemple:

Jemand hat eine $N(U) = 10'000$ mal zufällig eine Münze geworfen. $N(E) = 5094$ mal ist Kopf resp. Wappen realisiert worden. Da $h(E) = \frac{N(E)}{N(U)} = 0.5094$ gilt, wird $H_0 = (p_E = 0.5)$ angezweifelt. H_0 soll daher hier gegen die vermutete einseitige Alternative $H_1 = (p_E > 0.5)$ auf dem Signifikanzniveau $\alpha = 0.05$ getestet werden.

• Quelqu'un a jeté au hasard une monnaie $N(U) = 10'000$ fois. Il a réalisé $N(E) = 5094$ fois la face resp. l'armoire. Comme il vaut $h(E) = \frac{N(E)}{N(U)} = 0.5094$, on doute de l'hypothèse nulle $H_0 = (p_E = 0.5)$. Par conséquent il faut donc tester H_0 contre l'alternative unilatérale et supposée $H_1 = (p_E > 0.5)$ au niveau de signification $\alpha = 0.05$.

Sei X = Anzahl Kopf. Der Sollwert für $p = 0.5$ ist $X = 5'000$. Die Stichprobe ist verteilt nach der Binomialverteilung $Bi(10'000; 0.5) \hat{=} N(\mu, \sigma^2)$ (Die Binomialverteilung nähert sich der Normalverteilung an!). Dabei gilt:

• Soit X = nombre de faces. La valeur de consigne pour $p = 0.5$ est $X = 5'000$ (la distribution binomiale s'approche de la distribution normale). L'échantillon est distribué d'après la loi binomiale $Bi(10'000; 0.5) \hat{=} N(\mu, \sigma^2)$. Donc il vaut:

$$\mu = np, \quad \sigma^2 = npq = np(1-p), \quad P(X \leq c) = \frac{1}{\sqrt{2\pi}\sigma^2} \int_{-\infty}^x e^{-\frac{(u-\mu)^2}{2\sigma^2}} du,$$

$$d = \frac{c - \mu}{\sigma}, \quad Z = \frac{X - \mu}{\sigma} \Rightarrow P(X \leq c) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^d e^{-\frac{u^2}{2}} du = P(Z \leq d)$$

$$H_1 = (p_E > 0.5) \Rightarrow P(X > c)_{H_0=(p_E=0.5)} = \alpha \Rightarrow P(X \leq c) = P(Z \leq d) = 1 - \alpha = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^d e^{-\frac{u^2}{2}} du$$

In Zahlen unter der Voraussetzung der Nullhypothese $p = 0.5$:

• Avec des nombres sur la base de l'hypothèse nulle $p = 0.5$:

$$\mu = np = 10'000 \cdot 0.5 = 5'000, \quad \sigma^2 = npq = 10'000 \cdot 0.5 \cdot (1 - 0.5) = 10'000 \cdot 0.25 = 2'500, \quad \sigma = 50$$

$$Z = \frac{X - \mu}{\sigma} = \frac{X - 5'000}{50} = 0.1 X - 500, \quad X = Z \cdot 50 + 5'000 = 50 Z + 5'000$$

$$P(X \leq c) = P(Z \leq d) = 1 - \alpha = 0.95 = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^d e^{-\frac{u^2}{2}} du$$

Mathematica-Programm: • **Programme en Mathematica:**

```
alpha=0.05;
root = FindRoot[
  Integrate[E^(-u^2/2)/Sqrt[2 Pi], {u, -Infinity, d}] == 1-alpha, {d, 1}];
{d1 = d /. root, 50 d1 + 5000}
```

Output: • **Output:**

```
{1.64485, 5082.24}
```

Wir finden daher als Resultat die Schranke $c = 5082.24$ auf der Grundlage $\alpha = 0.05$. Die Wahrscheinlichkeit der Alternativhypothese $H_1 = (p_E > 0.5)$ ist $P(X > c)_{H_0=(p_E=0.5)} = P(X > 5082.24) = 0.05$. Dagegen haben wir aber festgestellt: $H(E) = N(E) = 5094$ bei $P(X > 5082.24) = 0.05$. Mit $H(E) = 5094$ ist daher so gesehen ein seltenes Ereignis mit der Wahrscheinlichkeit < 0.05 eingetreten. Daher kann man die Alternativhypothese nicht für zufällig halten und nimmt sie auf der Grundlage von $\alpha = 0.05$ an. Damit ist die Nullhypothese verworfen. Die Münze wird daher auf dieser Grundlage nicht für homogen gehalten. Wie man sieht, ist hier der Ablehnungsbereich oder kritische Bereich einseitig. ($\{X \mid X > 5082.24\}$).

• *Par conséquent nous trouvons comme le résultat la barrière à $c = 5082.24$ sur la base $\alpha = 0.05$. La probabilité de l'hypothèse alternative $H_1 = (p_E > 0.5)$ est $P(X > c)_{H_0=(p_E=0.5)} = P(X > 5082.24) = 0.05$. Mais par contre à cela nous avons constaté: $H(E) = N(E) = 5094$ avec $P(X > 5082.24) = 0.05$. Avec $H(E) = 5094$ il s'est donc réalisé un événement classifié comme rare avec la probabilité < 0.05 . Par conséquent on ne peut pas considérer l'hypothèse alternative comme accidentelle et on l'accepte donc sur la base de $\alpha = 0.05$. L'hypothèse nulle est donc repoussée. Par conséquent sur cette base la monnaie n'est pas considérée comme homogène. Comme on voit, le domaine de rejet ou **domaine critique** est ici unilatéral. ($\{X \mid X > 5082.24\}$).*

Direkte Rechnung — Calcul direct

Eine andere, kürzere Methode ist es, $p_1 = P(X > 5094)$ direkt zu berechnen und mit der signifikanten Wahrscheinlichkeit $\alpha = 0.05 = P(X > c)_{H_0=(p_E=0.5)}$ der Alternativhypothese $H_1 = (p_E > 0.5)$ zu vergleichen. Ist $p_1 < \alpha = 0.05$, so ist die Alternativhypothese daher nicht zufällig und die Nullhypothese daher abzulehnen.

• *Une autre méthode plus courte consiste à calculer directement $p_1 = P(X > 5094)$ et de comparer ce nombre avec la probabilité significative $\alpha = 0.05 = P(X > c)_{H_0=(p_E=0.5)}$ de l'hypothèse alternative $H_1 = (p_E > 0.5)$. Si on trouve $p_1 < \alpha = 0.05$, l'hypothèse alternative n'est pas aléatoire et il faut donc rejeter l'hypothèse nulle.*

Mathematica-Programm: • **Programme en Mathematica:**

```
sigma = 50; mu = 5000; cH = 5094;
p1 = 1/(Sqrt[2Pi ] sigma) Integrate[
  E^(-(mu-u)^2/(2 sigma^2)),{u, a, b}]/. {a->cH, b->Infinity} // N
```

Output: • **Output:**

```
| 0.030054 |
```

Konsequenz: • **Conséquence:** $p_1 = 0.030054 < \alpha = 0.5 \leadsto H_0$ ablehnen • *refuser H_0*

5.6.4 Testrisiken — Risques (aléas) aux tests

Aufgrund eines Tests könnte es möglich sein, dass man eine richtige Nullhypothese zufällig und fälschlicherweise ablehnt oder eine falsche Nullhypothese zufällig und fälschlicherweise nicht ablehnt. Etwas Ähnliches ist uns schon bei den Konfidenzintervallen begegnet. Wir übertragen daher die Begriffe von damals:

• *En raison d'un test, il peut être possible qu'on refuse accidentellement et fausement une hypothèse nulle correcte ou qu'on ne refuse pas accidentellement et fausement une fausse hypothèse nulle. Nous avons déjà rencontré une chose pareille aux intervalles de confiance. Par conséquent nous reprenons les notations:*

Definition: • **Définition:** Lehnen wir aufgrund eines Tests eine wahre Nullhypothese zufällig ab, so begehen wir einen **Fehler erster Art**.
 • *Si nous refusons par hasard une hypothèse nulle vraie en raison d'un test, nous commettons un **erreur de première sorte**.*

Einen Fehler 1. Art begehen wir, wenn wir die Alternativhypothese an Stelle der Nullhypothese annehmen. Die Alternativhypothese ist mit der Wahrscheinlichkeit α erfüllt (Signifikanzzahl).

• *Nous faisons une erreur de 1ère sorte si nous acceptons l'hypothèse alternative à la place de l'hypothèse nulle. L'hypothèse alternative est réalisée par la probabilité α (nombre significatif).*

Konsequenz: • **Conséquence:**

Die Wahrscheinlichkeit eines Fehlers 1. Art ist gerade die Signifikanzzahl α .
 • *La probabilité d'une erreur de 1ère sorte est exactement le nombre significatif α .*

Definition: • **Définition:** Lehnen wir aufgrund eines Tests eine falsche Nullhypothese zufällig nicht ab, so begehen wir einen **Fehler zweiter Art**.
 • *Si par hasard nous ne refusons pas une hypothèse nulle fausse en raison d'un test, nous commettons un **erreur de deuxième sorte**.*

Begehen wir einen Fehler zweiter Art, so wäre z.B. an Stelle der Nullhypothese $H_0 = (z = z_0)$ eine Alternativhypothese $H_1 = (z = z_1)$ richtig. Aufgrund von H_0 hat man eine Testfunktion konstruiert, die unrichtig ist. Die richtige Testfunktion müsste man aufgrund von H_1 konstruieren. Die Wahrscheinlichkeit für das Eintreffen der zu H_1 gehörigen Alternativhypothese bezüglich des aufgrund von H_0 berechneten c sei β . β ist somit gleich der Signifikanzzahl einer neuen Alternativhypothese bezüglich der neuen Nullhypothese H_1 . Somit ist $1 - \beta$ die Wahrscheinlichkeit die neue Nullhypothese H_1 nicht abzulehnen

und daher die alte Nullhypothese H_0 abzulehnen, d.h. einen Fehler 2. Art zu vermeiden.

• Si nous commettons une erreur de la deuxième sorte, à la place de l'hypothèse nulle, p.ex. $H_0 = (z = z_0)$, une hypothèse alternative $H_1 = (z = z_1)$ serait correcte. En raison de H_0 on a construit une fonction de test qui est fautive. On devrait construire la fonction de test correcte en raison de H_1 . β soit la probabilité pour la réalisation de l'hypothèse alternative liée à H_1 par rapport à c qui était calculé à la base de H_0 . Par conséquent β est égal au nombre significatif d'une nouvelle hypothèse alternative par rapport à la nouvelle hypothèse nulle H_1 . Par conséquent $1 - \beta$ est la probabilité de ne pas refuser la nouvelle hypothèse nulle H_1 et par conséquent de refuser la vieille hypothèse nulle H_0 , c.-à.-d. d'éviter une erreur de la 2ème sorte.

Definition: • **Définition:**

Die Wahrscheinlichkeit β , einen Fehler 2. Art zu vermeiden heisst **Macht des Tests**.

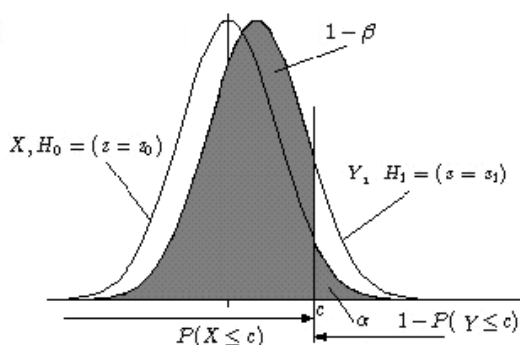
Die Wahrscheinlichkeit $1 - \beta$ einen Fehler 2. Art nicht zu vermeiden, heisst **Risiko 2. Art** oder **Produzentenrisiko**.

Die Signifikanzzahl oder Wahrscheinlichkeit eines Fehlers 1. Art heisst auch **Risiko 1. Art** oder **Konsumentenrisiko**.

• La probabilité β d'éviter une erreur de la 2ème sorte s'appelle le **pouvoir du test**.

La probabilité $1 - \beta$ de ne pas éviter une erreur la 2ème sorte s'appelle le **risque de la 2ème sorte** ou le **risque du producteur**.

Le nombre significatif ou la probabilité d'une erreur de la 1ère sorte s'appelle aussi le **risque de la 1ère sorte** ou le **risque du consommateur**.



$$\alpha = P(c < X)_{H_0}, \quad \beta = 1 - P(\tilde{X} \leq c)_{H_1}, \\ H_1 = (u = u_0) \Rightarrow \beta = \beta(u)$$

Bemerkung: • **Remarque:**

Ein grösseres c bedeutet zwar ein kleineres α , jedoch auch ein grösseres $1 - \beta$. Mit grösseren n bekommt man aber in der Regel schlankere Kurven. Bei gegebenem α wird dadurch $1 - \beta$ kleiner. Der Test gewinnt an **Trennschärfe**.

• Un c plus grand signifie en effet un α plus petit, cependant ça signifie aussi un $1 - \beta$ plus grand. Avec un n plus grand on obtient normalement des courbes plus fines. α étant donné, $1 - \beta$ devient plus petit. Le **degré de séparation** du test augmente.

5.6.5 Chi-quadrat-Test für die Varianz — Test khi deux pour la variance

Bsp.: • **Exemple:**

Sei • Soit $X_i \in N(\mu, \sigma^2)$, $\sigma_0 = 50$, $n = 20$, $s^2 = 61.3$, $H_0 = (\sigma^2 = \sigma_0^2)$, $H_1 = (\sigma^2 > \sigma_0^2)$

Bemerkung: • **Remarque:**

$\sigma^2 > \sigma_0^2$ ist die einzige interessante Alternative.

• $\sigma^2 > \sigma_0^2$ est la seule alternative intéressante.

$$\frac{S^2}{\sigma_0^2} = \frac{1}{n-1} \sum_{i_1}^n \frac{(X_i - \bar{X})^2}{\sigma^2} = \frac{1}{n-1} \sum_{i_1}^n Z_i^2, \quad Y_i \in N(0, 1) \Rightarrow T = \frac{S^2}{\sigma_0^2} (n-1) \text{ hat eine } \chi_{n-1}^2\text{-Verteilung}$$

$$\bullet T = \frac{S^2}{\sigma_0^2} (n-1) \text{ est réparti d'après la loi } \chi_{n-1}^2.$$

Sei $\bullet$ Soit $\alpha = 0.05 \rightsquigarrow P(T > c)_{\sigma^2=50} = \alpha = 0.05 \Rightarrow P(T \leq c)_{\sigma^2=50} = 1 - \alpha = 0.95$

Mathematica-Programm: $\bullet$ **Programme en Mathematica:**

```
chi[x_, n_] := 1/(2^(n/2) Gamma[n/2]) x^((n - 2)/2) E^(-x/2);
alpha = 0.05; sQ = 61.3; sigmQ = 50; n = 20;
{"c", c = c /. FindRoot[
  1 - alpha == Evaluate[Integrate[chi[u, n - 1], {u, 0, c}]], {c, 40}]
// Chop, "Interv", {xU = 0, sQc = c sigmQ/(n - 1) }}
```

Output: $\bullet$ **Output:**

```
| {"c", 30.1435, "Interv", {0, 79.325}} |
```

Die Alternativhypothese besagt, dass ein Wert mit der Wahrscheinlichkeit $\alpha = 0.05$ ausserhalb des Intervalls $\{0, 79.325\}$ zu liegen kommt. $s^2 = 61.3$ liegt jedoch innerhalb des Intervalls. Es ist demnach kein Grund vorhanden die Alternativhypothese anzunehmen und damit die Nullhypothese zu verwerfen.

$\bullet$ *L'hypothèse alternative signifie qu'une valeur quelconque soit située à l'extérieur de l'intervalle $\{0, 79.325\}$ avec une probabilité $\alpha = 0.05$. $s^2 = 61.3$ cependant est situé à l'intérieur de l'intervalle. Il n'existe donc pas de raison d'accepter l'hypothèse alternative et de repousser l'hypothèse nulle.*

5.6.6 Automatisches testen — Tester automatiquement

Bsp.: $\bullet$ **Exemple:**

Mit Mathematica, Zitat aus den Unterlagen: (Man beachte die vorausgesetzte Normalverteilung und die t -Verteilung!) $\bullet$ *Avec Mathematica, citation de la littérature officielle: (Il faut considérer la base des lois normales et de Student!)*

A test of a statistical hypothesis is a test of assumption about the distribution of a variable. Given sample data, you test whether the population from which the sample came has a certain characteristic. You can use functions in this package to test hypotheses concerning the mean, the variance, the difference in two population means, or the ratio of their variances.

The data that is given as an argument in a test function is assumed to be normally distributed. As a consequence of the Central Limit Theorem, you can disregard this normality assumption in the case of tests for the mean when the sample size, n , is large and the data is unimodal. The test functions accept as arguments the list of univariate data, a hypothesized parameter and relevant options.

Hypothesis tests for the mean are based on the normal distribution when the population variance is known, and the Student t distribution with a degrees of freedom when the variance has to be estimated. If you know the standard deviation instead of the variance, you can also specify KnownStandardDeviation $\rightarrow$ std.

The output of a hypothesis test is a pvalue, which is the probability of the sample estimate being as extreme as it is given that the hypothesized population parameter is true. A twosided test can be

requested using `TwoSided → True`. For more detailed information about a test use `FullReport → True`. This causes the parameter estimate and the test statistic to be included in the output. You can also specify a significance level using `SignificanceLevel → siglev`, which yields a conclusion of the test, stating acceptance or rejection of the hypothesis.

Mathematica-Programm: • **Programme en Mathematica:**

```
<< Statistics`HypothesisTests`;
data = {35, 33, 42, 32, 42, 43, 37, 44, 41, 39};
MeanTest[data, 34, KnownVariance -> 8, SignificanceLevel -> 0.9]
```

Output: • **Output:**

```
{OneSidedPValue -> 4.01256 * 10^-8,
 "Reject null hypothesis at significance level" -> 0.9}
```

Mathematica-Programm: • **Programme en Mathematica:**

```
MeanTest[data1, 38, KnownVariance -> 10, SignificanceLevel -> 0.1]
```

Output: • **Output:**

```
{OneSidedPValue -> 0.211855,
 "Fail to reject null hypothesis at significance level" -> 0.1}
```

Mathematica-Programm: • **Programme en Mathematica:**

```
MeanTest[data1, 38, KnownVariance -> 10]
```

Output: • **Output:**

```
{OneSidedPValue -> 0.211855,
 OneSidedPValue -> 0.211855}
```

Mathematica-Programm: • **Programme en Mathematica:**

```
MeanTest[data1, 38]
```

Output: • **Output:**

```
| OneSidedPValue -> 0.286061 |
```

Mathematica-Programm: • Programme en Mathematica:

```
| MeanTest[data1, 38, KnownVariance -> 10, FullReport -> True] // TeXForm |
```

Output: • Output:

```
| \{ {\Mvariable{FullReport}}\rightarrow  
  {\matrix{ 38.8 & 0.8 & \Muserfunction{NormalDistribu  
            tion}() \cr } },  
  {\Muserfunction{OneSidedPValue}}\rightarrow  
  {0.211855}}\}
```

↪

```
{FullReport→ Mean      TestStat      Distribution  
 38.8          0.8          NormalDistribution[], OneSidedPValue→0.211855}
```

5.6.7 Testrezept im Überblick — Vue d'ensemble d'une recette de test

1. Fasse das zu untersuchende Problem in eine Hypothese $\rightsquigarrow$ Nullhypothese z.B. $H_0 = (a = a_0)$. Gegeben sind hier ein Wert a_0 und ein Schätzer $a = \tilde{a}_0$. Zum Wert a soll eine Verteilungsfunktion X existieren mit speziell $X = a$.
 - *Résumer le problème en question par une hypothèse $\rightsquigarrow$ hypothèse nulle p.ex. $H_0 = (a = a_0)$. Ici on a donné une valeur a_0 et une estimation $a = \tilde{a}_0$. A la valeur a il doit exister une fonction de répartition X avec spécialement $X = a$*
2. Über $H_0 = (a = a_0)$ ist oft direkt keine Wahrscheinlichkeitsaussage möglich. Formuliere daher zu H_0 eine Alternativhypothese H_1 , über die mittels einer geeigneten Verteilungsfunktion und einem Intervall eine Wahrscheinlichkeitsaussage gemacht werden kann. Ein Beispiel ist die einseitige Alternativhypothese $H_1 = (a > a_0)$.
 - *Souvent il n'est pas directement possible de donner un énoncé (assertion) de probabilité concernant $H_0 = (a = a_0)$. Par rapport à H_0 on formule donc une hypothèse alternative H_1 dont il est possible de trouver une assertion de probabilité à l'aide d'une fonction de répartition et d'un intervalle convenable. Un exemple est l'hypothèse alternative unilatérale $H_1 = (a > a_0)$.*
3. Transformiere X bijektiv und stetig in eine andere Zufallsvariable $Z = g(X)$ (Testverteilung!), deren Verteilfunktion direkt berechenbar ist oder in Form von Tabellen vorliegt. $\rightsquigarrow X = g^{-1}(Z)$. Wenn g bijektiv und stetig ist, sind g und g^{-1} streng monoton. Ungleichungen bezüglich Z transformieren sich daher in Ungleichungen bezüglich X .
 - *Transformer X de façon bijective et monotone en une autre variable aléatoire $Z = g(X)$ (répartition de test!), dont la fonction de répartition est directement calculable ou dont on a donné les valeurs dans un tableau. $\rightsquigarrow X = g^{-1}(Z)$. Si g est bijective et continue, g et g^{-1} sont strictement monotones. Des inégalités par rapport à Z se transforment en inégalités par rapport à X .*
4. Zu $H_1 = (a > a_0) = H_1 = (a_0 < a)$ ist bei einem geeigneten a_1 aufgrund der Verteilung von X dann die Aussage " $a_0 < a \leq a_1$ " oder " $a_0 < X \leq a_1$ " wahrscheinlich und " $a_1 < a$ " oder " $a_1 < X$ " unwahrscheinlich. Transformiert haben diese Aussagen wegen der Monotonie z.B. die

Form " $g(a_0) < g(X) = Z \leq g(a_1)$ " resp. z.B. zu " $g(a_1) = c < g(X) = Z$ ".

- Par rapport à $H_1 = (a > a_0) = H_1 = (a_0 < a)$ et en considération de la répartition de X , la proportion " $a_0 < a \leq a_1$ " ou " $a_0 < X \leq a_1$ " a une probabilité grande si a_1 est une valeur convenable que la proportion " $a_1 < a$ " ou " $a_1 < X$ " soit improbable. A cause de la monotonie, ces proportions obtiennent après la transformation p.ex. la forme " $g(a_0) < g(X) = Z \leq g(a_1)$ " resp. p.ex. la forme " $g(a_1) = c < g(X) = Z$ ".

5. Zu einer vorgeschlagenen Wahrscheinlichkeit (Signifikanzzahl α) kann man nun die Ungleichung $P(c < Z) = \alpha$ resp. $P(Z \leq c) = 1 - \alpha$ mit Computern oder Tabellen nach c auflösen. $\leadsto c$
 - Soit donné une probabilité proposée (nombre significatif α); on peut résoudre l'inéquation $P(c < Z) = \alpha$ resp. $P(Z \leq c) = 1 - \alpha$ d'après c à l'aide d'ordinateurs ou de tableaux. $\leadsto c$
6. $c < Z$ entspricht z.B. $a_1 = g^{-1}(c) < X = g^{-1}(Z)$. Ist nun der Schätzer $\tilde{a}_0 \in (c, \infty)$, so hat man ein seltenes Ereignis der Alternativhypothese in einem Intervall mit kleiner Wahrscheinlichkeit α . Die Realität „fällt also in ein unwahrscheinliches Intervall“! Man kann daher die Alternativhypothese nicht verwerfen. Daher lehnt man in diesem Fall die Nullhypothese ab.
 - $c < Z$ correspond p.ex. à $a_1 = g^{-1}(c) < X = g^{-1}(Z)$. S'il vaut maintenant pour l'estimation $\tilde{a}_0 \in (c, \infty)$, un a un événement rare pour l'hypothèse alternative dans un intervalle où la probabilité α est petite. La réalité "tombe dans un intervalle improbable"! On ne peut donc pas rejeter l'hypothèse alternative. C'est pourquoi on rejette ici l'hypothèse nulle.

5.6.8 Vergleichstest für Mittelwerte I — Test de compar. de moyennes I

Man kann auf den ersten Blick zwei verschiedenartige Stichproben unterscheiden: Erstens solche, die aus verschiedenen Individuen bestehen und zweitens solche, die aus verschiedenen gewonnenen Messwerten an denselben Individuen bestehen. Wir betrachten hier einmal den zweiten Fall.

- On voit toute suite qu'on peut distinguer deux échantillons divers: Premièrement ceux qui consistent en individus différents et deuxièmement ceux qui consistent en valeurs obtenues de façon différentes, mais obtenues aux mêmes individus. Nous considérons ici le deuxième cas.

Bsp.: • **Exemple:** **Geg.:** • **Donné:**

Zwei Stichproben: Paarweise am selben Werkstück an zwei verschiedenen Orten A , B mit dem Verfahren V_A gemessene Daten und mit dem Verfahren V_B gemessene Daten:

- Deux échantillons: On a mesuré à deux endroits différents A , B à la même pièce des données par la méthode V_A et par la méthode V_B :

x_i	y_i	$d_i = x_i - y_i$
15.46	14.37	+0.09
16.35	16.36	-0.01
14.11	14.03	+0.08
19.21	19.23	-0.02
17.89	17.82	+0.07
15.66	15.62	+0.04

Wir haben es hier mit zwei verbundenen Stichproben zu tun mit paarweise zusammengehörigen Werten (x_i, y_i) und den zugehörigen Variablen (X_i, Y_i) . Wir wollen voraussetzen, dass alle X_i resp. Y_i unabhängig von i sind und gleiche Varianzen σ_X^2 resp. σ_Y^2 haben.

- Nous avons ici à faire à deux échantillons conjoints qui contiennent des valeurs en paires (x_i, y_i) et les variables affiliées (X_i, Y_i) . Nous voulons présupposer que tous les X_i resp. Y_i soient indépendants de i et aient des variances σ_X^2 resp. σ_Y^2 égales.

Unter diesen Voraussetzungen wollen wir die Hypothesen $H_0 = (\mu_{X_i} = \mu_{Y_i})$ gegen die einseitigen Alternativen $H_1 = (\mu_{X_i} > \mu_{Y_i})$ testen.

• *A ces conditions, nous voulons tester les hypothèses $H_0 = (\mu_{X_i} = \mu_{Y_i})$ contre les alternatives unilatérales $H_1 = (\mu_{X_i} > \mu_{Y_i})$.*

Wir wissen, dass die Variablen Gaussverteilungen haben und daher zu den Differenzen $d_i = x_i - y_i$ eine gaussverteilte Variable D gehört mit $\mu_D = \mu_X - \mu_Y$ und $\sigma_D^2 = \sigma_X^2 + \sigma_Y^2$. Nun können wir mit Hilfe des t -Tests die Nullhypothese $\mu_D = 0$ gegen die Alternative $\mu_D > 0$ testen.

• *Nous savons que les variables sont distribuées selon la loi de Gauss et que par conséquent les différences $d_i = x_i - y_i$ ont une variable D qui est distribuée selon la loi de Gauss avec $\mu_D = \mu_X - \mu_Y$ et $\sigma_D^2 = \sigma_X^2 + \sigma_Y^2$. Maintenant nous pouvons tester à l'aide du test de Student l'hypothèse nulle $\mu_D = 0$ contre l'alternative $\mu_D > 0$.*

$$\leadsto \bar{d} = \frac{1}{n} \sum_{i_1}^n d_i = 0.0416667, \quad s_D = \sqrt{\frac{1}{n-1} \sum_{i_1}^n (d_i - \bar{d})^2} = 0.0470815, \quad \mu_0 = 0, \quad \alpha = 0.05, \quad n = 6$$

$$Z = g(X) = T = \frac{\bar{D} - \mu_0}{S/\sqrt{n}} = \frac{X - \mu_0}{S/\sqrt{n}} \leadsto X = \bar{D} = g^{-1}(Z) = \mu_0 + \frac{T S}{\sqrt{n}} = \mu_0 + \frac{Z S}{\sqrt{n}}$$

$$P(c < T = Z) = \alpha = 0.05 \Rightarrow P(T \leq 0) = 1 - \alpha \leadsto c$$

$$c < Z \Leftrightarrow g^{-1}(c) < g^{-1}(Z) \Leftrightarrow \mu_0 + \frac{c S}{\sqrt{n}} = g^{-1}(c) < g^{-1}(Z) = X$$

$$\Leftrightarrow 0 + c \frac{0.0470815}{\sqrt{6}} = c \cdot 0.0192209 = g^{-1}(c) < g^{-1}(Z) = X$$

$\leadsto$ Frage: • *Question: $c \cdot 0.0192209 \stackrel{?}{<} X = \bar{d} = 0.0416667$* $\leadsto$

Mathematica-Programm: • **Programme en Mathematica:**

```
alpha = 0.05;
f[z_, m_] :=
  Gamma[(m + 1)/2]/(Sqrt[m Pi] Gamma[m/2]) / (1 + z^2/m)^((m + 1)/2);
FindRoot[
  1 - alpha == Evaluate[Integrate[f[t, 5], {t, -Infinity, c}]], {c, 3}
// Chop
```

Output: • **Output:**

```
{c -> 2.01505}
```

$$\leadsto c \cdot 0.0192209 = 2.01505 \cdot 0.0192209 = 0.0387311 < 0.0416667 = \bar{d}$$

$\leadsto$ Seltenes Ereignis, Alternativhypothese nicht verwerfen, Nullhypothese verwerfen.

• *Événement rare, ne pas rejeter l'hypothèse alternative, rejeter l'hypothèse nulle.*

5.6.9 Vergleichstest für Mittelwerte II — Test de compar. de moyennes II

Wir betrachten den Fall, dass zwei unabhängige Stichproben $\{x_1, \dots, x_n\}$ für $X \in N(\mu_X, \sigma_X^2)$ und $\{y_1, \dots, y_m\}$ für $Y \in N(\mu_Y, \sigma_Y^2)$ mit allgemein $n = n_X \neq m = n_Y$ und $\sigma_X = \sigma_Y$ vorliegen. Gegeben seien die Schätzungen $\bar{x}$ für μ_X und entsprechend $\bar{y}$, s_X , s_Y . Getestet werden soll die Nullhypothese $H_0 = (\mu_X = \mu_Y)$ gegen die Alternative $H_0 = (\mu_X > \mu_Y)$. Man sieht rasch, dass man wegen $n_X \neq n_Y$ jetzt nicht mehr mit den Differenzen eine neue Variable konstruieren kann. Um trotzdem eine Testvariable zu konstruieren, braucht es hier also etwas mehr Theorie. Des Rahmens wegen verzichten wir hier auf eine weitere Herleitung und begnügen uns mit dem Resultat:

• *Nous considérons le cas de deux échantillons indépendants $\{x_1, \dots, x_n\}$ pour $X \in N(\mu_X, \sigma_X^2)$ et $\{y_1, \dots, y_m\}$ pour $Y \in N(\mu_Y, \sigma_Y^2)$ avec en générale $n = n_X \neq m = n_Y$ et $\sigma_X = \sigma_Y$. Soient données les estimations $\bar{x}$ pour μ_X et correspondamment $\bar{y}$, s_X , s_Y . Il faut tester l'hypothèse nulle $H_0 = (\mu_X = \mu_Y)$ par contre à l'alternative $H_0 = (\mu_X > \mu_Y)$. On voit très vite que maintenant à cause de $n_X \neq n_Y$ on ne peut plus construire avec les différences une nouvelle variable. Pour construire quand même une variable de test, on a donc besoin d'un peu plus de théorie. À cause du cadre, nous renonçons ici à une dérivation plus vaste et nous nous contentons du le résultat:*

Formel: • **Formule:**

Die folgende Variable genügt unter unseren Voraussetzungen einer t -Verteilung mit $n_X + n_Y - 2$ Freiheitsgraden:

• *A nos conditions, la variable suivante satisfait la loi de Student avec $n_X + n_Y - 2$ degrés de liberté :*

$$T_{XY} = \sqrt{\frac{n_X n_Y (n_X + n_Y - 2)}{n_X + n_Y}} \cdot \frac{\bar{X} - \bar{Y}}{\sqrt{(n_X - 1) S_X^2 + (n_Y - 1) S_Y^2}}$$

Bsp.: • **Exemple:**

$\alpha = 0.05$

x_i	y_i	$d_i = x_i - y_i$
15.46	14.37	+0.09
16.35	16.36	-0.01
14.11	14.03	+0.08
19.21	19.23	-0.02
17.89	17.82	+0.07
15.66	15.62	+0.04
15.48		

Sei • *Soit* $H_0 = (\mu_X = \mu_Y)$, $H_1 = (\mu_X > \mu_Y)$

↪ **Mit Mathematica:** • **A l'aide de Mathematica:**

a = {15.46, 16.35, 14.11, 19.21, 17.89, 15.66, 15.48};

b = {14.37, 16.36, 14.03, 19.23, 17.82, 15.62};

Length[a]

7

Length[b]

6

```
LocationReport[a]
```

```
{Mean -> 16.3086, HarmonicMean -> 16.1616, Median -> 15.66}
```

```
LocationReport[b]
```

```
{Mean -> 16.2383, HarmonicMean -> 16.0372, Median -> 15.99}
```

```
DispersionReport[a]
```

```
{Variance -> 2.93031, StandardDeviation -> 1.71182, SampleRange -> 5.1,
  MeanDeviation -> 1.29265, MedianDeviation -> 0.69,
  QuartileDeviation -> 1.02}
```

```
DispersionReport[b]
```

```
{Variance -> 4.04326, StandardDeviation -> 2.01079, SampleRange -> 5.2,
  MeanDeviation -> 1.565, MedianDeviation -> 1.725,
  QuartileDeviation -> 1.725}
```

```
alpha = 0.05; n = Length[a] + Length[b] - 2;
```

```
f[z_, m_] :=
```

```
  Gamma[(m + 1)/2]/(Sqrt[m Pi] Gamma[m/2]) / (1 + z^2/m)^((m + 1)/2);
```

```
c = c /.
```

```
  FindRoot[
```

```
    1 - alpha == Evaluate[Integrate[f[u, n], {u, -Infinity, c}]], {c, 3}]
```

```
// Chop
```

```
1.79588
```

```
nX = Length[a]; nY = Length[b];
```

```
mX = Mean /. LocationReport[a][[1]];
```

```
mY = Mean /. LocationReport[b][[1]];
```

```
sX = StandardDeviation /. DispersionReport[a][[2]];
```

```
sY = StandardDeviation /. DispersionReport[b][[2]];
```

```
{nX, nY, mX, mY, sX, sY}
```

```
{7, 6, 16.3086, 16.2383, 1.71182, 2.01079}
```

```
const = Sqrt[((nX - 1)sX^2 + (nY - 1)sY^2)(nX + nY)/(nX nY (nX + nY - 1))]
```

```
0.987397
```

```
c * const
```

```
1.77325
```

```
c * (mX - mY)
```

```
0.12614
```

$$\leadsto P(c < T_{XY}) = \alpha, \quad T = T_{XY} = \sqrt{\frac{n_X n_Y (n_X + n_Y - 2)}{n_X + n_Y}} \cdot \frac{\bar{X} - \bar{Y}}{\sqrt{(n_X - 1) S_X^2 + (n_Y - 1) S_Y^2}}$$

$$\leadsto (c < T_{XY})_{\alpha=0.05} \Leftrightarrow |_{\alpha=0.05} :$$

$$c \cdot \sqrt{\frac{((n_X - 1) S_X^2 + (n_Y - 1) S_Y^2) \cdot (n_X + n_Y)}{n_X n_Y (n_X + n_Y - 2)}} < (\bar{x} - \bar{x}) \cdot \sqrt{\frac{((n_X - 1) S_X^2 + (n_Y - 1) S_Y^2) \cdot (n_X + n_Y)}{n_X n_Y (n_X + n_Y - 2)}}$$

$$\leadsto 1.77325 < 0.12614$$

Die letzte Aussage trifft aber nicht zu, d.h. $(c < T_{XY})$ ist hier nicht erfüllt. Die Alternativhypothese kann man also nicht annehmen und die Nullhypothese daher nicht ablehnen.

• *La dernière proposition n'est pas vraie, c.-à-d. $(c < T_{XY})$ n'est pas satisfait ici. On ne peut donc pas accepter l'hypothèse alternative et par conséquent on ne peut pas refuser l'hypothèse nulle.*

5.7 Weitere Testverfahren — Autres méthodes de test

5.7.1 Eine Typeneinteilung von Tests — Une classification des tests

Grundsätzlich kann man sich beliebige Kriterien ausdenken, um Tests in Klassen einzuteilen. Für uns leisten die folgenden zwei Typologien ihre Dienste:

• *En principe on peut inventer des critères quelconques pour classifier les tests. Pour nous les deux typologies suivantes nous rendent service:*

1. **Parametertests**, z.B. verteilungsabhängige Tests. Z.B. $H_0 = (\mu_X = \mu_Y)$ soll getestet werden.
 - *Des tests de paramètres, p.ex. il faut tester $H_0 = (\mu_X = \mu_Y)$.*
 2. **Parameterfreie Tests**, z.B. die **Rangtests** oder der **Vorzeichentest**.
 - *Des tests libres de paramètres, p.ex. les tests de rang ou le test des signes (signum).*
-

1. **Verteilungsabhängige Tests**, z.B. Parametertests.
 - *Les tests qui dépendent d'une loi de répartition, p.ex. des tests de paramètres.*
 2. **Verteilungsunabhängige Tests**, z.B. Rangtests. Oder auch **Verteilungs- oder Anpassungstests**, mit denen geprüft wird, ob eine vermutete Verteilung zu den empirischen Werten passt.
 - *Les tests qui ne dépendent pas d'une loi de répartition, p.ex. des tests de rang. Ou en outre des tests de la qualité de la loi de distribution où on examine si une répartition est compatible avec les valeurs empiriques.*
-

5.7.2 Parameterfreier Test: Der Vorzeichentest — Test libre de paramètres: Le test du signe

Bsp.: • Exemple: Geg.: • Donné:

Tägliche Umsatzmengen von zwei gleichen Arbeitsteams A und B mit identischen Aufträgen während 10 Tagen. Notiert ist nur die Differenz $W_A - W_B$ in Verrechnungseinheiten. Man teste die Nullhypothese,

dass beide Teams gleich gut sind gegen die Alternative, dass man mit Team A einen systematisch höheren Umsatz erzielt als mit B . Sei $\alpha = 0.05$ (klein): Wenn die Realität für H_1 eine Wahrscheinlichkeit $P < \alpha$ ergibt, ist H_1 bemerkenswert, unwahrscheinlich oder „signifikant“.

• *Quantités de vente quotidiennes de deux équipes de travail comparables A et B avec des tâches identiques pendant 10 jours. On a noté seulement la différence $W_A - W_B$ en unités de calcul. On teste l'hypothèse nulle où les deux équipes ont les mêmes capacités contre l'alternative où l'équipe A réalise une vente systématiquement plus haute que B. Soit $\alpha = 0.05$ (petit): Si de la réalité on obtient comme résultat une probabilité $P < \alpha$, H_1 devient remarquable, improbable ou "significative".*

$W_A - W_B$	0.9	2.1	1.1	0.0	1.1	-0.4	1.4	0.6	0.5	0.1
signum	+	+	+	·	+	-	+	+	+	+

Es gilt: • *Il vaut:* $H_0 = (W_A = W_B)$ wahr • *vrai* $\leadsto P(\text{sgn}(W_A - W_B) = " + ") = P(\text{sgn}(W_B - W_A) = " + ")$

Der Wert 0.0 kann gestrichen werden, da er die Entscheidung nicht ändern kann. Somit haben wir 8 mal " + " unter 9 Werten.

• *La valeur 0.0 peut être tracée, car la décision ne peut pas changer. Ainsi nous avons 8 fois " + " parmi 9 valeurs.*

Idee: • **Idée:** Sei $X = \text{Anzahl " + "}$ unter den n Werten. Unter der Voraussetzung von H_0 ist X binomialverteilt mit $p = 0.5$ ($P(\text{sgn}(W_A - W_B) = " + ") = P(\text{sgn}(W_B - W_A) = " + ")$).

• *Soit $X = \text{nombre de " + "}$ entre les n valeurs. Sous la condition de H_0 X est réparti selon la loi binomiale avec $p = 0.5$ ($P(\text{sgn}(W_A - W_B) = " + ") = P(\text{sgn}(W_B - W_A) = " + ")$).*

$$\begin{aligned} \leadsto P(X \geq 8) &= 1 - P(X \leq 7) = 1 - \sum_{k=0}^7 P(X = k) = 1 - \sum_{k=0}^7 \binom{9}{k} p^k q^{9-k} \\ &= 1 - \sum_{k=0}^7 \binom{9}{k} 0.5^k 0.5^{9-k} = 1 - 0.5^9 \sum_{k=0}^7 \binom{9}{k} = 1 - 0.5^9 (2^9 - \binom{9}{8} - \binom{9}{9}) \approx 0.0195313 \\ &\Rightarrow P(X \geq 8) \approx 0.0195313 \leq \alpha = 0.05 \end{aligned}$$

Man hat es also hier mit einem seltenen Ereignis mit $X = 8$ zu tun, das eingetroffen ist, obwohl $P(X \geq 8) \approx 0.0195313 \leq P(X \geq c) = \alpha = 0.05$ gilt. Aus denselben Überlegungen wie vorher mussten wir daher eine Alternativhypothese akzeptieren und die Nullhypothese daher verwerfen. Wäre hingegen $\alpha = 0.01$, so könnte man die Alternativhypothese nicht akzeptieren und die Nullhypothese daher nicht verwerfen.

• *On l'a donc à faire ici à un événement rare avec $X = 8$ qui a été réalisé, bien qu'il vaille $P(X \geq 8) \approx 0.0195313 \leq P(X \geq c) = \alpha = 0.05$. Pour les mêmes réflexions qu'avant, il faut donc accepter une hypothèse alternative et rejeter par conséquent l'hypothèse nulle. Mais si l'on avait $\alpha = 0.01$, on ne pourrait pas accepter d'hypothèse alternative et on ne pourrait donc plus rejeter l'hypothèse nulle.*

5.7.3 Chi-Quadrat-Anpassungstest — Test d'adaptation de Khi deux

Problem: • **Problème:** Es ist mittels eines Tests zu prüfen, ob eine bestimmte Variable eine gegebene Verteilungsfunktion $F(t)$ besitzt. Wir wollen das an einem Beispiel studieren:

• *Il faut examiner au moyen d'un test, si une variable donnée suit la loi d'une fonction de distribution $F(t)$ donnée. C'est ce que nous voulons étudier par l'exemple suivant:*

Bsp.: • **Exemple:** **Geg.:** • **Donné:**

1. Disjunkte Ereignisse $A_1, \dots, A_m$ als Resultate eines Zufallsexperiments. • *Événements disjoints $A_1, \dots, A_m$ qui sont des résultats d'une expérience aléatoire.*

$$A_i \cap A_k = \{\}, \quad i \neq k, \quad \bigcup_{i=1}^m = \Omega$$

2. $P(A_i) = p_i$ unbekannt • *inconnu* $\leadsto H_0 = (P(A_i) = p_i), i = 1, \dots, n$
3. Test von H_0 : Nehme Stichprobe vom Umfang $n, n_i = |A_i|$
 • *Test de H_0 : Prendre un échantillon de la taille $n, n_i = |A_i|$*
4. Bilde: • *Formule:* $\tau = \sum_{i=1}^m \frac{(n p_i - n_i)^2}{n p_i} = \sum_{i=1}^m \frac{(\nu_i - n_i)^2}{\nu_i} = \sum_{i=1}^m \nu_i \left(1 - \frac{n_i}{\nu_i}\right)^2$

Es lässt sich zeigen, dass τ für grosse m annähernd χ^2 verteilt ist mit $m - 1$ Freiheitsgraden, vgl. z.B. Lit. Kreyszig, Bibl. A10 sowie Bibl. A7 (Cramer).

• *on peut démontrer que τ est une approximation de la loi de χ^2 avec $m - 1$ degrés de liberté, si m est assez grand, voir. p.ex. lit. Kreyszig, Bibl. A10 et Bibl. A7 (Cramer).*

Entscheidungsfindung: Analoge Überlegungen wie schon bei den obigen Beispielen führen uns zum Entscheid H_0 abzulehnen, sobald der Wert von τ grösser ist als $c = \chi_{\alpha, m-1}^2$ (c = obere Grenze des Integrals mit dem Wert α).

• **Trouver la décision:** *Des réflexions analogiques aux exemples qu'on vient de traiter nous mènent à la décision de refuser H_0 dès que la valeur de τ surmonte $c = \chi_{\alpha, m-1}^2$ (c = borne supérieure de l'intégrale à la valeur α).*

Bsp.: • **Exemple:** (vgl. Lit Kreyszig, Bibl. A10.)

Gregor Mendel, Kreuzungsversuche mit Erbsen. p_i vermutet nach den „Mendelschen Gesetzen“ .

• *Gregor Mendel, expériences avec des pois, hybrides. p_i supposé d'après les "lois de Mendel".*

Vermutung: • *Supposition:* $n_A : n_B : n_C : n_D = 9 : 3 : 3 : 1 \leadsto$ Resultat: • *Résultat:*

	A	B	C	D	$m = 4$
	rund, gelb • <i>rond, jaune</i>	rund, grün • <i>rond, vert</i>	kantig, gelb • <i>arêts, jaune</i>	kantig, grün • <i>arêts, vert</i>	Total • <i>total</i>
n_i	315	108	101	32	556
p_i	$\frac{9}{16}$	$\frac{3}{16}$	$\frac{3}{16}$	$\frac{1}{16}$	1
$n p_i$	312.75	104.25	104.25	34.75	556

Test der Vermutung mit $\alpha = 0.05$:

• *Test de la supposition avec $\alpha = 0.05$: $n_A : n_B : n_C : n_D = 9 : 3 : 3 : 1$*

$$\tau = \frac{(312.75 - 315)^2}{312.75} + \frac{(104.25 - 108)^2}{104.25} + \frac{(104.25 - 101)^2}{104.25} + \frac{(34.75 - 32)^2}{34.75}$$

$\leadsto$ **Mit Mathematica:** • **A l'aide de Mathematica:**

```
(312.75 - 315)^2/312.75 + (104.25 - 108)^2/104.25 + (104.25 - 101)^2/
104.25 + (34.75 - 32)^2/34.75
```

0.470024

```
chi[x_, n_] := 1/(2^(n/2) Gamma[n/2]) x^((n - 2)/2) E^(-x/2);
alpha = 0.05; m = 4; {"c",
  c = c /. FindRoot[
    1 - alpha == Evaluate[Integrate[chi[u, m - 1], {u, 0, c}]], {c,
    8}] // Chop}
```

$\{ "c", 7.81471 \}$

$\leadsto \leadsto$ Entscheid: • *Décision:* $\tau = 0.470024 < c_{\alpha=0.5} = 7.81471 \leadsto H_0$ nicht abgelehnt! Problem: τ ist für grosse m annähernd χ^2 verteilt. Und bei uns war $m = 4 \dots$ • H_0 *n'est pas refusée!* *Problème:* τ est une approximation de la loi de χ^2 , si m est assez grand. Et chez nous on a $m = 4 \dots$

Bemerkung: • **Remarque:**

Ein Ereignis A_i könnte es auch sein, dass X in ein gewisses Intervall $[x_i, x_{i+1})$ fällt. Mit Intervallen ist auch die kontinuierliche Situation vertreten. So lässt sich z.B. **testen, ob eine Funktion normalverteilt ist.**

• *Un événement A_i pourrait aussi être que X tombe dans un certain intervalle $[x_i, x_{i+1})$. Avec les intervalles, la situation continue est aussi représentée. Ainsi on peut tester si une fonction est distribuée de façon normale.*

5.7.4 Zum Fishertest — Quant au test de Fisher

Auf Seite 148 haben wir gesehen: • *A la page 148 vous avons vu:*

Satz: • **Théorème:** **Vor.:** • **Hyp.:** $\mathcal{F}$ -Verteilung • *Distribution $\mathcal{F}$*
 Beh.: • **Thè.:**

$$\begin{aligned}
 1. \quad F(x) = P(\mathcal{F} \leq x) &= \frac{\Gamma(\frac{m_1 + m_2}{2})}{\Gamma(\frac{m_1}{2})\Gamma(\frac{m_2}{2})} \cdot m_1^{\frac{m_1}{2}} \cdot m_2^{\frac{m_2}{2}} \cdot \int_0^x \frac{t^{\frac{m_1}{2}-1}}{(m_1 t + m_2)^{\frac{m_1 + m_2}{2}}} dt \\
 2. \quad f(x) &= \begin{cases} 0 & x \leq 0 \\ f(x) = E_{m_1, m_2} x^{\frac{m_1}{2}-1} (m_2 + m_1 \cdot x)^{-\frac{m_1 + m_2}{2}} & x > 0 \end{cases}
 \end{aligned}$$

Nun gilt (Beweise vgl. Lit. Kreyszig, Bibl. A10):

• *Maintenant il vaut (preuves voir lit. Kreyszig, Bibl. A10):*

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Zufallsvariablen $V_1, V_2 \in \{\heartsuit\}$ sowie $V_1 \in \chi_{m_1}^2$, $V_2 \in \chi_{m_2}^2$
 • *Variables aléatoires* $V_1, V_2 \in \{\heartsuit\}$ ainsi que $V_1 \in \chi_{m_1}^2$, $V_2 \in \chi_{m_2}^2$

Beh.: • **Thè.:**

$V = \frac{V_1/m_1}{V_2/m_2}$ besitzt eine $\mathcal{F}$ -Verteilung • *est réparti selon la loi $\mathcal{F}$*

$$x \leq 0 \Rightarrow F(x) = 0, \quad 0 < x \Rightarrow F(x) = P(V \leq x) = \\ = \frac{\Gamma(\frac{m_1+m_2}{2})}{\Gamma(\frac{m_1}{2})\Gamma(\frac{m_2}{2})} \cdot m_1^{\frac{m_1}{2}} \cdot m_2^{\frac{m_2}{2}} \cdot \int_0^x \frac{t^{\frac{m_1}{2}-2}}{(m_1 t + m_2)^{\frac{m_1+m_2}{2}}} dt$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Zufallsvariablen • *Variables aléatoires*
 $X_1, \dots, X_{n_1}, Y_1, \dots, Y_{n_2} \in \{\heartsuit\}$ sowie $X_i \in N(\mu_1, \sigma_1^2)$, $Y_i \in N(\mu_2, \sigma_2^2)$
 S_1^2, S_2^2 Schätzer für • *Estimations de σ_1^2, σ_2^2* (Empirische Varianzen)
 • *Variances empiriques*

Beh.: • **Thè.:**

$V = \frac{S_1^2/\sigma_1^2}{S_2^2/\sigma_2^2}$ besitzt eine $\mathcal{F}$ -Verteilung mit $(n_1 - 1, n_2 - 1)$ Freiheitsgraden
 • *est réparti selon la loi $\mathcal{F}$ avec $(n_1 - 1, n_2 - 1)$ degrés de liberté*

Konsequenz: • **Conséquence:**

Bei $H_0 = (\sigma_1^2 = \sigma_2^2)$ wird $V = \frac{S_1^2}{S_2^2}$. Damit hat man ein Mittel um die Gleichheit der Varianzen zu testen!

• *Pour $H_0 = (\sigma_1^2 = \sigma_2^2)$ on obtient $V = \frac{S_1^2}{S_2^2}$. On a donc un moyen pour tester l'égalité des variances!*

5.7.5 Kontingenztafeln — Tableaux de contingence

Mit dem χ^2 -Test kann man auch die **statistische Unabhängigkeit** zweier Zufallsvariablen X und Y testen. Wir betrachten dazu ein Beispiel einer Kontingenztafel. Darin sind Kunden, die entweder das Produkt A , B oder C kaufen, in Altersklassen eingeteilt. Aufgeführt sind die Häufigkeiten (empirische Werte):

• *A l'aide du test de χ^2 on peut aussi tester l'indépendance statistique de deux variables aléatoires X et Y . Nous considérons un exemple de tableaux de contingence. Les clients qui achètent le produit A , B ou C y sont répartis dans des classes d'après leur âge. On y voit les fréquences (valeurs empiriques):*

n_{ik}	$Y = A$	$Y = B$	$Y = C$	$n_{i\cdot}$
$X \in [16, 30)$	$n_{11} = 132$	$n_{12} = 778$	$n_{13} = 592$	$n_{1\cdot} = 1502$
$X \in [30, 55)$	$n_{21} = 55$	$n_{22} = 304$	$n_{23} = 248$	$n_{2\cdot} = 607$
$X \in [55, 85)$	$n_{31} = 25$	$n_{32} = 111$	$n_{33} = 155$	$n_{3\cdot} = 291$
$n_{\cdot k}$	$n_{\cdot 1} = 212$	$n_{\cdot 2} = 1193$	$n_{\cdot 3} = 995$	$n = 2400$

$n_{\cdot k}$, $n_{i\cdot} \rightsquigarrow$ Randhäufigkeiten • *fréquence marginales*

Problem: • **Problème:** Sind X und Y statistisch unabhängig? • *Est-ce-que X et Y sont statistiquement indépendants?*

Idee: • **Idée:**

Die empirischen Werte $n_{ik}, n_{i\cdot}, n_{\cdot k}$ betrachten wir als Realisationen von zu erwartenden theoretischen Besetzungszahlen $n \cdot p_{ik}, n \cdot p_{i\cdot}, n \cdot p_{\cdot k}$. Die gestellte Frage können wir nun in eine Nullhypothese fassen. Statistische Unabhängigkeit bedeutet, dass für die Wahrscheinlichkeiten gelten muss: $p_{ik} = p_{i\cdot} \cdot p_{\cdot k}$.

• *Nous considérons les valeurs empiriques $n_{ik}, n_{i\cdot}, n_{\cdot k}$ comme réalisations de nombres d'occupation théoriques et attendues $n \cdot p_{ik}, n \cdot p_{i\cdot}, n \cdot p_{\cdot k}$. Nous pouvons maintenant poser la question comme hypothèse nulle. L'indépendance statistique signifie que pour la probabilité il vaut: $p_{ik} = p_{i\cdot} \cdot p_{\cdot k}$.*

$$\leadsto H_0 = (p_{ik} = p_{i\cdot} \cdot p_{\cdot k})$$

$\frac{n_{ik}}{n}$ ist ein Schätzwert für p_{ik} , $\frac{n_{i\cdot}}{n}$ ist ein Schätzwert für $p_{i\cdot}$, $\frac{n_{\cdot k}}{n}$ ist ein Schätzwert für $p_{\cdot k}$. Unabhängigkeit müsste also bedeuten:

• $\frac{n_{ik}}{n}$ est une valeur estimée pour p_{ik} , $\frac{n_{i\cdot}}{n}$ est une valeur estimée pour $p_{i\cdot}$, $\frac{n_{\cdot k}}{n}$ est une valeur estimée pour $p_{\cdot k}$. L'indépendance doit donc signifier:

$$\leadsto H_0 = (p_{ik} = p_{i\cdot} \cdot p_{\cdot k}) \approx H_0 = \left(\frac{n_{ik}}{n} = \frac{n_{i\cdot}}{n} \cdot \frac{n_{\cdot k}}{n}\right) = H_0 = \left(n_{ik} = \frac{n_{i\cdot} \cdot n_{\cdot k}}{n}\right)$$

Weiter kann man zeigen (vgl. Lit.): • *En outre on peut démontrer:*

Satz: • **Théorème:**

Die folgende Testvariable ist approximativ $\chi^2_{(r-1) \cdot (s-1)}$ -verteilt:

• *La variable aléatoire suivante est distribuée approximativement selon la loi $\chi^2_{(r-1) \cdot (s-1)}$:*

$$T = \sum_{i=1}^r \sum_{k=1}^s \frac{\left(\frac{n_{i\cdot} \cdot n_{\cdot k}}{n} - n_{ik}\right)^2}{\frac{n_{i\cdot} \cdot n_{\cdot k}}{n}}$$

Mittels der Randhäufigkeiten berechnen wir nun die geschätzten erwarteten Besetzungszahlen:

• *Au moyen des fréquences aux limites nous calculons les valeurs d'occupation estimées attendues:*

$$H_0 \leadsto n \cdot p_{ik} \approx \frac{n_{i\cdot} \cdot n_{\cdot k}}{n} = u_{ik}$$

u_{ik}	$Y = A$	$Y = B$	$Y = C$
$X \in [16, 30)$	132.677	746.619	622.704
$X \in [30, 55)$	53.6183	301.73	251.652
$X \in [55, 85)$	25.705	144.651	120.644

$$\text{Es gilt: } \bullet \text{ Il vaut: } T = \sum_{i=1}^r \sum_{k=1}^s \frac{\left(\frac{n_{i\cdot} \cdot n_{\cdot k}}{n} - n_{ik}\right)^2}{\frac{n_{i\cdot} \cdot n_{\cdot k}}{n}} = \sum_{i=1}^r \sum_{k=1}^s \frac{(u_{ik} - n_{ik})^2}{u_{ik}} \Rightarrow T = t = \dots$$

Sei • *Soit* $\alpha = 0.01$

$\leadsto$ **Mit Mathematica:** • **A l'aide de Mathematica:**

```
v1 = {{212, 1193, 995}}; v2 = {{1502, 607, 291}};
matrV = {{132, 778, 592}, {55, 304, 248}, {25, 111, 155}};
matrU = Transpose[v2].v1/2400 //
```

```

N; matrV = {{132, 778, 592}, {55, 304, 248}, {25, 111, 155}};
t={1, 1, 1}.({1, 1, 1}.((matrU - matrV)^2/matrU)) // Evaluate

20.5737

chi[x_, n_] := 1/(2^(n/2) Gamma[n/2]) x^((n - 2)/2)E^(-x/2);
alpha = 0.01; r = 3; s = 3; m = (r - 1)(s - 1);
{"c", c = c /. FindRoot[
  1 - alpha == Evaluate[Integrate[chi[u, m], {u, 0, c}]], {c, 8}] //
Chop}

{"c", 13.2766}

```

→ Entscheid: • *Décision*: $t = 20.5737 \stackrel{?}{<} c_{\alpha=0.1} = 13.2766$

→ H_0 abgelehnt! Problem: T ist für grosse r, s annähernd χ^2 verteilt. Und bei uns war $r = s = 4 \dots$

• H_0 refusée! Problème: T est une approximation de la loi de χ^2 , si r, s sont assez grands. Et chez nous on a $r = s = 4 \dots$

Kapitel 6

Fehlerrechnung, Regression, Korrelation — Calcul de l'erreur, régression, corrélation

Auszug aus: • *Extrait de:* Wirz, Bibl. A15, Script $\diamond$ Math $\diamond$ Ing, $\diamond$ Analysis $\diamond$ Analyse $\diamond$, kurz & bündig. Mit Ergänzung • *Avec supplément*

6.1 Fehlerrechnung (Abhängigkeit) — Calcul de l'erreur (dépendance)

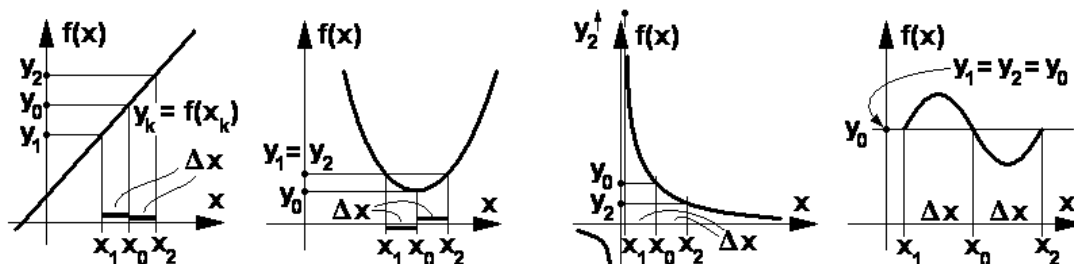
Unter dem Begriff „Fehlerrechnung“ werden in der Literatur zwei verschiedene Problemtypen behandelt: Erstens das Problem der Verpflanzung von Messfehlern bei der Anwendung von Funktionen auf fehlerbehaftete Messgrößen. Und zweitens das Problem der Fehler von statistischen Kenngrößen (z.B. Mittelwert), welche meistens von einer grossen Datenmenge abhängen. Hier wollen wir den ersten Fehler-typ behandeln. Der zweite wird im getrennt herausgegebenen Anhang zu diesem Skript besprochen.

• *Sous la notion de calcul d'erreur, deux types de problèmes différents sont traités dans la littérature: Premièrement le problème de la transplantation d'erreurs de mesure à l'application de fonctions lors de grandeurs mesurées inexactes. Et deuxièmement le problème des erreurs de grandeurs qui marquent des données statistiques (p.ex. la moyenne) qui dépend d'ordinaire d'une grande quantité de données. Ici, nous voulons traiter le premier type d'erreur. Le deuxième est traité dans l'appendice (édition séparée) à ce script.*

6.1.1 Das Problem der Verpflanzung — Le problème de la dépendance

Situation: • **Situation:** (a): $y = f(x)$

Gemessen wird • *On mesure* $x \pm \Delta x \rightsquigarrow y \pm \Delta y = ?$



$$y_0 = f(x_0), \quad y_1 = f(x_1), \quad y_2 = f(x_2), \quad \Delta y_1 = |y_1 - y_0| = ?, \quad \Delta y_2 = |y_2 - y_0| = ?$$

Situation: • **Situation:** (b): $y = f(x_{1_0}, x_{2_0}, \dots, x_{n_0})$

Gemessen werden die Werte $x_{1_0}, x_{2_0}, \dots, x_{n_0}$ der Variablen $x_1, x_2, \dots, x_n$. Bei kontinuierlichen Messwerten gibt es immer Ablesungenauigkeiten, die aber abschätzbar sind. Diese zugehörigen „Messfehler“ betragen $\Delta x_1, \Delta x_2, \dots, \Delta x_n$. Die „exakten Werte“ x_k^* , $k = 1, \dots, n$ liegen daher in den Intervallen $[x_{k_0} - \Delta x_k, x_{k_0} + \Delta x_k]$. Zudem sei eine Funktion $f(x_1, x_2, \dots, x_n)$ gegeben, mit deren Hilfe eine weitere Grösse berechnet werden muss. • *On mesure les valeurs $x_{1_0}, x_{2_0}, \dots, x_{n_0}$ des variables $x_1, x_2, \dots, x_n$. Pour les valeurs de mesures non-discrètes il y a toujours des inexactitudes quand on lit les échelles, mais elles sont estimables. Les erreurs de mesure qui en résultent sont $\Delta x_1, \Delta x_2, \dots, \Delta x_n$. Les "valeurs exactes" x_k^* , $k = 1, \dots, n$ sont donc situées dans les intervalles $[x_{k_0} - \Delta x_k, x_{k_0} + \Delta x_k]$. En plus on a donné une fonction $f(x_1, x_2, \dots, x_n)$ à l'aide de laquelle on doit calculer une valeur en plus.*

Problem: • **Problème:** In welchem Intervall liegt der „wahre“ Wert $f(x_1^*, x_2^*, \dots, x_n^*)$?

• *Dans quel intervalle la valeur "exacte" $f(x_1^*, x_2^*, \dots, x_n^*)$ est-elle située ?*

6.1.2 Verwendung des totalen Differentials — Appliquer la différentielle totale

Sei • *Soit* $\vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$, $\vec{x}_0 = \begin{pmatrix} x_{1_0} \\ \vdots \\ x_{n_0} \end{pmatrix}$, $f(\vec{x}) := f(x_1, x_2, \dots, x_n)$, $D_k \geq |\Delta x_k|$ (D_k ist eine bezifferbare Schranke. • *D_k est une borne connue.*)

Aus der Theorie des **totalen Differentials** weiss man: • *De la théorie de la différentielle totale on sait:*

$$\Delta f = f(\vec{x}_0 + \Delta \vec{x}) - f(\vec{x}_0) = \Delta x_1 f'_{x_1}(\vec{x}_0) + \dots + \Delta x_n f'_{x_n}(\vec{x}_0) + O[2]$$

(O: Glieder höherer Ordnung • *Termes d'ordre supérieur*)

$$\leadsto \Delta f \approx \Delta x_1 f'_{x_1}(\vec{x}_0) + \dots + \Delta x_n f'_{x_n}(\vec{x}_0)$$

$$\leadsto |\Delta f| \leq |\Delta x_1| |f'_{x_1}(\vec{x}_0)| + \dots + |\Delta x_n| |f'_{x_n}(\vec{x}_0)| \leq D_1 |f'_{x_1}(\vec{x}_0)| + \dots + D_n |f'_{x_n}(\vec{x}_0)| := \Delta f_{\max}$$

Satz: • **Théorème:** **Vor.:** • **Hyp.:**

Messsituation wie oben beschrieben • *Situation de mesure comme décrit en haut,*
 $f \in \mathcal{D}^{(1)}$

Beh.: • **Thè.:**

$$|\Delta f| \leq D_1 |f'_{x_1}(\vec{x}_0)| + \dots + D_n |f'_{x_n}(\vec{x}_0)| := \Delta f_{\max}$$

Konsequenz: • **Conséquence:**

$$f(\vec{x}^*) = f(x_1^*, x_2^*, \dots, x_n^*) \in [f(\vec{x}_0) - \Delta f_{\max}, f(\vec{x}_0) + \Delta f_{\max}]$$

Definition: • **Définition:** $|\Delta f|$ heisst **absoluter Fehler** • *s'appelle erreur absolue*,
 $|\frac{\Delta f}{f(\vec{x}_0)}|$ heisst **relativer Fehler** • *s'appelle erreur relative*.

1. Beispiel: • **Exemple 1:** $f(x, y) = x \pm y \Rightarrow \Delta f_{max} = D_x |1| + D_y |\pm 1| = D_x + D_y$

2. Beispiel: • **Exemple 2:** $f(x, y) = x \cdot y \Rightarrow \Delta f_{max} = D_x |y_0| + D_y |x_0|$

3. Beispiel: • **Exemple 3:** $f(x, y) = \frac{x}{y} \Rightarrow \Delta f_{max} = D_x |\frac{1}{y_0}| + D_y |\frac{x_0}{y_0^2}|$

4. Beispiel: • **Exemple 4:** $f(x, y) = x^y \Rightarrow \Delta f_{max} = D_x |y_0 \cdot x_0^{y_0-1}| + D_y |x_0^{y_0} \ln(x_0)|$

5. Beispiel: • **Exemple 5:**

$$f(x) = x^2 - 2x + 4 - \sin(x) + \ln(x) \Rightarrow \Delta f_{max} = D_x |2x_0 - 2 - \cos(x_0) + \frac{1}{x}|$$

Bemerkung: • **Remarque:** Diese Beispiele zeigen, dass die oft geäußerte Meinung, es genüge mit den extremen Werten zu rechnen, wohl äusserst falsch sein muss. • *Ces exemples démontrent que l'opinion souvent communiquée, qu'il suffit de calculer avec les valeurs extrêmes, doit être complètement fausse.*

6. Beispiel: • **Exemple 6:**

Messwerte: • *Valeurs de mesure:*

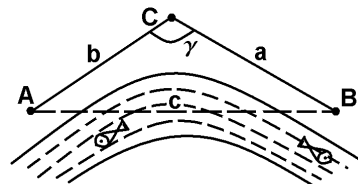
$$a = 364.76 \pm 0.05m$$

$$b = 402.35 \pm 0.05m$$

$$\gamma \hat{=} 68^\circ 14' \pm 4'$$

$$\leadsto \gamma \approx 1.1909 \pm 0.002$$

$$c = ?$$



$$\leadsto c = \sqrt{a^2 + b^2 - 2ab \cos(\gamma)} \approx 431.38$$

$$\Rightarrow \Delta c_{max} = D_a \cdot \left| \frac{\partial c}{\partial a} \right| + D_b \cdot \left| \frac{\partial c}{\partial b} \right| + D_\gamma \cdot \left| \frac{\partial c}{\partial \gamma} \right| =$$

$$= 0.05 \cdot \left| \frac{2a - 2b \cos(\gamma)}{2\sqrt{a^2 + b^2 - 2ab \cos(\gamma)}} \right| + 0.05 \cdot \left| \frac{2b - 2a \cos(\gamma)}{2\sqrt{a^2 + b^2 - 2ab \cos(\gamma)}} \right| + 0.002 \cdot \left| \frac{2ab \sin(\gamma)}{2\sqrt{a^2 + b^2 - 2ab \cos(\gamma)}} \right|$$

$$\approx 0.02498 + 0.03096 + \underbrace{0.36762}_{!!!} \approx 0.424 \Rightarrow c \pm \Delta c_{max} = 431.38 \pm 0.424$$

6.2 Regression — Régression

6.2.1 Der Begriff — La notion

Geg.: • **Donné:**

Stichprobe von Beobachtungen • *Epreuve au hasard d'observations* $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$

x_k z.B. Maschineneinstellung • *p.ex. réglage (positionnement) de la machine*,

$y_k = f(x_k)$ Messung • *measurement*.

Problem: • **Problème:** Interpretation des Verhaltens der Messwerte (Gesetz)?

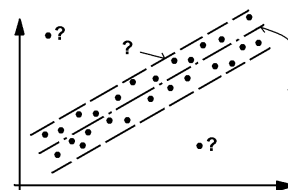
• *Intérpretation du comportement des mesures (lois)?*

Z.B. Vermutung: Die Messwerte liegen auf einer passenden Geraden oder sonstigen Kurve.

• *P.ex. supposition: Les points mesures sont situés sur une droite qui convient ou sur une autre courbe.*

Problem: • **Problème:**

Wie findet man die „beste“ Gerade (resp. Kurve)? • *Comment trouver la "meilleure" droite (resp. courbe)?*



~> **Probleme:** • **Problèmes:**

1. Berechnung der Parameter der hypothetischen Kurve. • *Calculer les paramètres de la courbe hypothétique.*
2. Beurteilung der Zuverlässigkeit der getroffenen Wahl. • *Juger l'authenticité du choix qu'on a fait.*

Die so gefundene Kurve trägt den etwas unverständlichen Namen **Regressionskurve**.

• *La courbe ainsi trouvée porte le nom un peu incompréhensible de courbe de régression.*

Wieso dieser Name? Der Name stammt aus einer Beschreibung von Beobachtungen von F. Galton²¹ zum Grössenwachstum von Menschen. Galton hat festgestellt: • *Pourquoi ce nom? Le nom vient d'une description d'observations de F. Galton²¹ quant à la croissance de l'être humain. Galton a observé:*

- 1 Grössere Väter haben grössere Söhne. • *Les pères plus grands ont des fils plus grands.*
- 2 Jedoch beobachtet man die Tendenz, dass grosse Väter grosse Söhne haben, die aber im Mittel kleiner sind als die genannten Väter selbst. Es ist also ein **Rückschritt** ~> **Regress** vorhanden.
• *Mais par contre on observe la tendance que les pères plus grands ont des fils grands mais qui sont en moyenne plus petits que les pères dont on parle. Il s'agit donc d'un régression.*

Sei • *Soit* x = Grösse der Väter • *grandeur des pères,*
 y = Grösse der Söhne • *grandeur des fils.*
 $y(x) = ax + b$ Gerade • *droite*
 ~> „**Regressionsgerade**“ • **„droite de régression“**

Der Begriff *Regressionsgerade* ist also historisch verankert, hat aber inhaltlich nichts mit der Mathematik zu tun. • *La notion de droite de régression est donc fondée par l'histoire. La signification de cette notion n'a rien à faire avec les mathématiques.*

6.2.2 Methode der kleinsten Quadrate — Méthode des carrés minimaux

Das Problem — Le problème

Problem: • **Problème:** Welche Kurve ist die „beste“ Kurve? – Wie ist das Kriterium „beste“ definitorisch am vernünftigsten zu fassen? • *Quelle est la „meilleure“ des courbes? – Comment est-ce qu'il faut définir le critère "meilleur" dans ce cadre de façon raisonnable?*

Da wir Menschen selbst entscheiden müssen, was unsere „Vernunft“ sein soll, können wir hier „vernünftig“ mit „ökonomisch“ und daher mit „pragmatisch“ und „einfach“ gleichsetzen. Wir wählen daher hier ein pragmatisches Vorgehen. (In der Literatur ist eine etwas fundiertere Begründung üblich, auf der Grundlage des „Maximum-likelihood-Prinzips.“) • *Comme nous, êtres humains, devons décider nous-mêmes ce qui est notre "raison", nous pouvons remplacer ici "raisonnable" par "économique", donc par "pragmatique" et "simple". C'est pourquoi nous choisissons ici un chemin pragmatique. (Il est de coutume dans la littérature de présenter une explication un peu plus étoffée, basée sur le principe de "maximum-likelihood".)*

²¹Engl. Naturforscher • *Naturaliste anglais*, 1822 – 1911

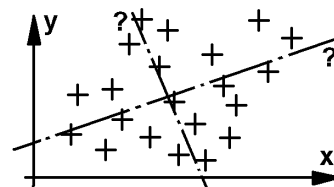
Wichtig: • Important: Dass die gesuchte Kurve eine Gerade oder sonst irgend eine Kurve ist, lässt sich theoretisch nicht exakt entscheiden. Man kann nur zu Aussagen kommen wie „die eine Kurve ist wahrscheinlicher als die andere“. Denn dass die gefundenen Messwerte überhaupt auf einer stetigen Kurve liegen, beruht auf einer Arbeitshypothese, die wir nach Descartes mit dem Argument der Erfahrung und der Arbeitsökonomie begründen. • *On ne peut pas décider de façon théorique que la courbe cherchée est une droite ou n'importe quelle autre courbe. On peut seulement affirmer que "l'une des courbes est plus propable que l'autre". C'est une hypothèse que les valeurs mesurées sont situées sur une courbe continue, hypothèse que nous faisons d'après Descartes, nous basant sur l'expérience et l'économie du travail.*

Die Begründung der Methode — La déduction de la méthode

Problem: • Problème: Durch eine „Punktwolke“ von Messwerten soll eine „beste“ Gerade gelegt werden. Wie ist vorzugehen? • *Il faut tracer la "meilleure" droite à travers un "nuage de points mesures". Comment faut-il procéder?*

Idee 1: • Idée 1:

Versuche g so festzulegen, dass die Summe der Abstände zu g gleich 0 wird. • *Essaie de déterminer g de façon que la somme des distances à g devienne 0.*

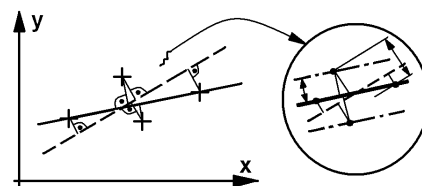


Problem: • Problème: Diese Methode funktioniert nicht, da offensichtlich (Skizze) mehrere passende „beste“ Geraden möglich sind. • *Cette méthode ne fonctionne pas car visiblement (esquisse) plusieurs droites du type "meilleure" sont possibles.*

Idee 2: • Idée 2:

Lege die Gerade so, dass die Summe der Beträge der Abstände minimal wird.

• *Choisis la droite de façon que la somme des valeurs absolues des distances à g soit minimale.*

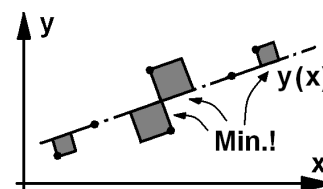


Problem: • Problème: Diese Methode funktioniert praktisch schlecht, da beim Berechnen der Abstände im $\mathbb{R}^2$ Wurzeln vorkommen, was die Rechnung sehr kompliziert. • *Cette méthode fonctionne mal en pratique, car en calculant les distances dans le $\mathbb{R}^2$ on obtient des racines carrées ce qui complique passablement le calcul.*

Idee 3: • Idée 3:

Nimm statt der Summe der Beträge der Abstände die Summe der Quadrate der Abstände zu g . Dadurch fallen bei der Rechnung die Quadratwurzeln weg.

• *Choisis au lieu de la somme des valeurs absolues des distances à g la somme des carrés des valeurs absolues des distances. Par conséquent on élimine de cette manière les racines carrées.*



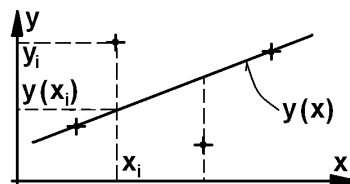
Problem: • Problème: Praktisch ist zu einem gegebenen Wert x_i der Messwert y_i gegeben $\leadsto P_i$. Die Gerade soll so bestimmt werden, dass für die nächstgelegenen Punkte $P_i^* \in g$ die Summe der

Distanzquadrate $|\overline{P_i P_i^*}|$ minimal ist. Man hat das Problem der Minimalisierung unter der Bedingung, dass $\overline{P_i P_i^*} \perp g$ gilt, was die Rechnung wieder sehr kompliziert. Einfacher wird es, hier einen Fehler in Kauf zu nehmen und $x_i^* = x_i$ zu setzen. • *En pratique on a pour une valeur x_i donnée une valeur de mesure y_i donnée $\leadsto P_i$. La droite doit être déterminée de manière que pour les points les plus proches $P_i^* \in g$ la somme des carrés de distance $|\overline{P_i P_i^*}|$ soit minimale. Le problème de la minimalisation est lié à la condition suivante: $\overline{P_i P_i^*} \perp g$. Cela complique le calcul. Il est plus simple d'accepter une erreur et de mettre $x_i^* = x_i$.*

Idee 4: • Idée 4:

Nimm statt Summe der Quadrate der Abstände zu g nur die Summe der Δy_i .

• *Choisis, au lieu de la somme des carrés des valeurs absolues des distances à g , seulement la somme des Δy_i .*



Problem: • Problème: Unter der Hypothese, dass die Kurve eine Gerade $y = g(x) = a \cdot x + b$ ist, gilt es nun, a und b zu berechnen. Analog geht man bei andern Kurven vor, z.B. $y = h(x) = a \sin(bx + c)$.

• *Sous l'hypothèse que la courbe soit une droite $y = g(x) = a \cdot x + b$, il faut maintenant calculer a et b . On procède de façon analogue pour d'autres courbes, p.ex. $y = h(x) = a \sin(bx + c)$ etc..*

Lösung: • Solution:

$$(1) \sum_{i=1}^n (\Delta y_i)^2 = \sum_{i=1}^n (y_i - g(x))^2 = \sum_{i=1}^n (y_i - (a x + b))^2 := f(a, b) \rightarrow \text{Min}$$

$$(2) \sum_{i=1}^n (\Delta y_i)^2 = \sum_{i=1}^n (y_i - h(x))^2 = \sum_{i=1}^n (y_i - (a \sin(bx + c)))^2 := f(a, b) \rightarrow \text{Min}$$

Bedingung: • *Condition:* $\frac{\partial f}{\partial a} = \frac{\partial f}{\partial b} = 0 \quad \leadsto$

$$\begin{aligned} \frac{\partial f}{\partial a} &= \sum_{i=1}^n 2(y_i - a x_i - b) \cdot (-x_i) = -2 \sum_{i=1}^n (y_i - a x_i - b) \cdot x_i = -2 \sum_{i=1}^n y_i \cdot x_i - a \sum_{i=1}^n x_i^2 - b \sum_{i=1}^n x_i \\ &= -2 \sum_{i=1}^n y_i \cdot x_i - a \sum_{i=1}^n x_i^2 - b \sum_{i=1}^n x_i = 0 \end{aligned}$$

$$\begin{aligned} \frac{\partial f}{\partial b} &= \sum_{i=1}^n 2(y_i - a x_i - b) \cdot (-1) = -2 \sum_{i=1}^n (y_i - a x_i - b) = 0 \\ &\Rightarrow \sum_{i=1}^n y_i - \sum_{i=1}^n a x_i - \sum_{i=1}^n b = \sum_{i=1}^n y_i - a \sum_{i=1}^n x_i - n b = 0 \end{aligned}$$

Sei • *Soit* $\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i$ (Mitterwert der x_i • *valeur moyenne des x_i*), $\bar{y} = \frac{1}{n} \sum_{i=1}^n y_i$

$$\begin{aligned} \Rightarrow & \sum_{i=1}^n x_i y_i = a \sum_{i=1}^n x_i^2 + b n \bar{x} \\ & \sum_{i=1}^n y_i = a \sum_{i=1}^n x_i + n b \Rightarrow y(\bar{x}) = a \bar{x} + b = \bar{y} \end{aligned}$$

$b = \bar{y} - a \bar{x}$ einsetzen: • *substituer:* ($\Rightarrow \bar{y} = a \bar{x} + b$)

$$\sum_{i=1}^n x_i y_i = a \sum_{i=1}^n x_i^2 + (\bar{y} - a \bar{x}) n \bar{x} = a \left(\sum_{i=1}^n x_i^2 - n \bar{x}^2 \right) + n \bar{x} \bar{y}$$

Hinweis: • Indication:

$$\begin{aligned} \text{Sei} \bullet \text{ Soit } \vec{d}_e &:= \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} := \vec{1}, \quad \vec{x} := \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix}, \quad \vec{y} := \begin{pmatrix} y_1 \\ y_2 \\ \vdots \\ y_n \end{pmatrix} \\ &\Rightarrow \sum_{i=1}^n x_i y_i = \langle \vec{x}, \vec{y} \rangle, \quad \sum_{i=1}^n x_i = \langle \vec{x}, \vec{d}_1 \rangle, \quad \sum_{i=1}^n y_i = \langle \vec{y}, \vec{d}_1 \rangle \quad \leadsto \end{aligned}$$

Formel: • **Formule:****Vor.:** • **Hyp.:**

Sei • Soit

 $y = a \cdot x + b$ Regressionsgerade g • *Droite de régression* g **Beh.:** • **Thè.:**

$$\begin{aligned} a &= \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} = \frac{\langle \vec{x}, \vec{y} \rangle - n \bar{x} \bar{y}}{\langle \vec{x}, \vec{x} \rangle - n \bar{x}^2} \\ b &= \bar{y} - \bar{x} \cdot \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} = \bar{y} - \bar{x} \cdot \frac{\langle \vec{x}, \vec{y} \rangle - n \bar{x} \bar{y}}{\langle \vec{x}, \vec{x} \rangle - n \bar{x}^2} \\ y(\bar{x}) &= a \bar{x} + b = \bar{y} \Rightarrow (\bar{x}, \bar{y}) \in g \end{aligned}$$

Schreibweise mit Varianz und Kovarianz — Façon d'écrire à l'aide de la variance et de la covarianceBetrachtung: • *Considération:*

$$\begin{aligned} s_x^2 &:= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \left(\sum_{i=1}^n (x_i^2 - 2x_i \bar{x} + \bar{x}^2) \right) = \frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - 2\bar{x} \sum_{i=1}^n x_i + \sum_{i=1}^n \bar{x}^2 \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - 2\bar{x} n \bar{x} + n \bar{x}^2 \right) = \frac{1}{n-1} \left(\left(\sum_{i=1}^n x_i^2 \right) - n \bar{x}^2 \right) = \frac{1}{n-1} \left(\left(\sum_{i=1}^n x_i^2 \right) - \frac{1}{n} \left(\sum_{i=1}^n x_i \right)^2 \right) \\ &= \frac{1}{n-1} (\langle \vec{x}, \vec{x} \rangle - n \bar{x}^2) \end{aligned}$$

In der mathematischen Statistik ist es üblich, die **Varianzen** s_x^2 , s_y^2 und die **Kovarianz** s_{xy} wie folgt zu definieren: • *Dans la statistique mathématique on a la coutume de définir les variances s_x^2 , s_y^2 et la covariance s_{xy} comme il suit:*

Definition: • **Définition:**

$$\begin{aligned} s_x^2 &:= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = \frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - \frac{1}{n} \left(\sum_{i=1}^n x_i \right)^2 \right) = \frac{1}{n-1} (\langle \vec{x}, \vec{x} \rangle - n \bar{x}^2) \\ s_y^2 &:= \frac{1}{n-1} \sum_{i=1}^n (y_i - \bar{y})^2 = \frac{1}{n-1} \left(\sum_{i=1}^n y_i^2 - \frac{1}{n} \left(\sum_{i=1}^n y_i \right)^2 \right) = \frac{1}{n-1} (\langle \vec{y}, \vec{y} \rangle - n \bar{y}^2) \\ s_{xy} &:= \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y}) = \frac{1}{n-1} \left(\left(\sum_{i=1}^n x_i y_i \right) - n \bar{x} \bar{y} \right) \\ &= \frac{1}{n-1} \left(\sum_{i=1}^n x_i y_i - \frac{1}{n} \left(\sum_{i=1}^n x_i \right) \left(\sum_{i=1}^n y_i \right) \right) = \frac{1}{n-1} (\langle \vec{x}, \vec{y} \rangle - n \bar{x} \bar{y}) \end{aligned}$$

Es gilt: • *Il vaut:* $\bar{x} = \frac{1}{n} \langle \vec{x}, \vec{1} \rangle$, $\bar{y} = \frac{1}{n} \langle \vec{y}, \vec{1} \rangle$.

Mit Hilfe der Varianz und der Kovarianz lässt sich die Regressionsgerade einfacher schreiben. Dabei benutzen wir: • *A l'aide de la variance et de la covariance on peut simplifier la formule pour la droite de régression. Pour cela nous utilisons:*

Lemma: • **Lemme:** $a = \frac{s_{xy}}{s_x^2}, \quad b = \bar{y} - \bar{x} \cdot \frac{s_{xy}}{s_x^2}$

(Durch ausmultiplizieren verifiziert man leicht: • *En multipliant les termes des deux côtés on vérifie facilement:* $s_{xy} = s_x^2 \cdot a$)

Korollar: • **Corollaire:**

Vor.: • **Hyp.:**

Sei • *Soit*

$y = a \cdot x + b$ Regressionsgerade • *Droite de régression*

Beh.: • **Thè.:**

$$y = \frac{s_{xy}}{s_x^2} \cdot x + \bar{y} - \bar{x} \cdot \frac{s_{xy}}{s_x^2}$$

Bemerkung: • **Remarque:**

Für die Summe der quadratischen Abstände gilt: • *Pour la somme des distances au carré il vaut:*

$$\begin{aligned} \sum_{i=1}^n (\Delta y_i)^2 &= \sum_{i=1}^n (y_i - a x_i - b)^2 = \sum_{i=1}^n (y_i - \bar{y} - a(x_i - \bar{x}))^2 \\ &= \sum_{i=1}^n (y_i - \bar{y})^2 - 2a \sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y}) + a^2 \sum_{i=1}^n (x_i - \bar{x})^2 \\ &= (n-1) \cdot (s_y^2 - 2a s_{xy} + a^2 s_x^2) = (n-1) \cdot (s_y^2 - 2a a s_x^2 + a^2 s_x^2) = (n-1) \cdot (s_y^2 - a^2 s_x^2) \end{aligned}$$

$$\leadsto \sum_{i=1}^n (\Delta y_i)^2 = 0 \Leftrightarrow s_y^2 = a^2 s_x^2 = \left(\frac{s_{xy}}{s_x^2}\right)^2 s_x^2 = \frac{s_{xy}^2}{s_x^2} \Leftrightarrow (s_x \cdot s_y)^2 = s_{xy}^2$$

Satz: • **Théorème:** $\sum_{i=1}^n (\Delta y_i)^2 = 0 \Leftrightarrow (s_x \cdot s_y)^2 = s_{xy}^2$

6.2.3 Korrelation — Corrélation

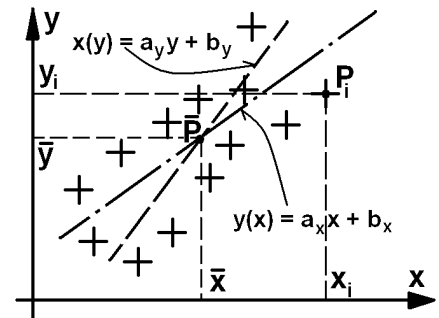
Vorhin haben wir die Gerade $g_x : y(x) = a \cdot x + b$ untersucht. • *Nous venons d'examiner la droite $g_x : y(x) = a \cdot x + b$.*

Da y die abhängige und x die unabhängige Variable war, schreiben wir präziser: • *Comme y était la variable dépendante et x la variable indépendante, nous écrivons de manière plus précise:*

$$y(x) = a_x \cdot x + b_x,$$

$$a_x = \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} = \frac{s_{xy}}{s_x^2} = \frac{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y})}{\frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2} = \frac{\sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y})}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Wir vertauschen nun alle Wertepaare (x_i, y_i) . D.h. wir betrachten y als die abhängige und x als die unabhängige Variable. Dann lässt sich mit der Methode der kleinsten Quadrate wieder die „beste“ Gerade $g_y : x(y) = a_y \cdot y + b_y$ durch die Messpunkte finden. • *Nous échangeons maintenant les paires de valeurs (x_i, y_i) . Ç.v.d. nous considérons y comme variable indépendante et x comme variable dépendante. Par conséquent on peut de nouveau trouver la droite la "meilleure" $g_y : x(y) = a_y \cdot y + b_y$ par les points de mesure à l'aide de la méthode des carrés minimaux.*



a_y und b_y erhält man aus a_x und b_x durch Rollenvertauschung der $x_i, \bar{x}$ und $y_i, \bar{y}$:

• *On obtient a_y et b_y de a_x et b_x par échangeement des rôles de $x_i, \bar{x}$ et $y_i, \bar{y}$:*

$$a_y = \frac{\sum_{i=1}^n y_i x_i - n \bar{y} \bar{x}}{\sum_{i=1}^n y_i^2 - n \bar{y}^2} = \frac{s_{yx}}{s_y^2} = \frac{\sum_{i=1}^n (y_i - \bar{y}) \cdot (x_i - \bar{x})}{\sum_{i=1}^n (y_i - \bar{y})^2},$$

$$b_y = \bar{x} - \bar{y} \cdot \frac{\sum_{i=1}^n y_i x_i - n \bar{y} \bar{x}}{\sum_{i=1}^n y_i^2 - n \bar{y}^2}$$

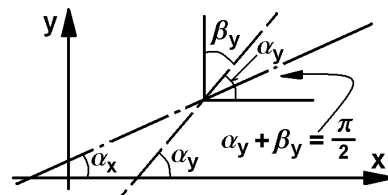
Wenn man g_x und g_y in dasselbe Koordinatensystem einzeichnet, ist zu erwarten, dass zwei verschiedene Geraden entstehen. • *Si on dessine g_x et g_y dans le même système de coordonnées, on peut s'attendre à deux droites différentes.*

Ideal wäre allerdings, wenn g_x und g_y zusammenfallen würden. Dann hätte man: • *Il serait idéal si les deux droites étaient les mêmes. Alors on aurait:*

$$a_x = \tan(\alpha_x),$$

$$a_y = \tan(\beta_y) = \tan\left(\frac{\pi}{2} - \alpha_x\right) = \frac{1}{\tan(\alpha_x)}$$

$$\Rightarrow a_x \cdot a_y = \tan(\alpha_x) \cdot \frac{1}{\tan(\alpha_x)} = 1$$

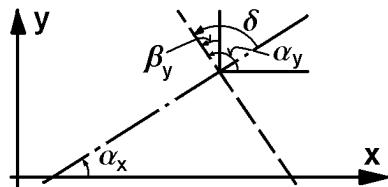


Andernfalls ist: • *Autrement on obtient:*

$$a_x \cdot a_y = \tan(\alpha_x) \cdot \tan(\beta_y)$$

$$= \frac{\tan(\alpha_x)}{\tan(\frac{\pi}{2} - \beta_y)} = \frac{\tan(\alpha_x)}{\tan(\alpha_y)}$$

$$= \frac{\tan(\alpha_x)}{\tan(\alpha_x + \delta)} \neq 1$$



$a_x \cdot a_y = \frac{\tan(\alpha_x)}{\tan(\alpha_x + \delta)}$ ist umso verschiedener von 1, je verschiedener δ von 0 ist • *est plus différent de 1, plus δ est différent de 0.* ~

$a_x \cdot a_y$ ist ein Mass für den Zusammenhang der beiden Geraden, d.h. für die **Korrelation**.

• *$a_x \cdot a_y$ est une mesure pour le rapport entre les deux droites, ç.v.d. pour la **corrélacion**.*

Es gilt: • *On constate:* $(s_x, s_y \geq 0)$

$$a_x \cdot a_y = \frac{\sum_{i=1}^n y_i x_i - n \bar{y} \bar{x}}{\sum_{i=1}^n x_i^2 - n \bar{x}^2} \cdot \frac{\sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}}{\sum_{i=1}^n y_i^2 - n \bar{y}^2} = \frac{(\sum_{i=1}^n y_i x_i - n \bar{y} \bar{x})^2}{(\sum_{i=1}^n x_i^2 - n \bar{x}^2) \cdot (\sum_{i=1}^n y_i^2 - n \bar{y}^2)},$$

$$\sqrt{a_x \cdot a_y \cdot \operatorname{sgn}(a_x \cdot a_y)} = \frac{|\sum_{i=1}^n y_i x_i - n \bar{y} \bar{x}|}{\sqrt{(\sum_{i=1}^n x_i^2 - n \bar{x}^2) \cdot (\sum_{i=1}^n y_i^2 - n \bar{y}^2)}} = \frac{|s_{xy}|}{\sqrt{s_x^2 \cdot s_y^2}} := |r_{xy}|$$

Definition: • **Définition:**

$r_{xy} = \frac{s_{xy}}{\sqrt{s_x^2 \cdot s_y^2}}$ heisst **Korrelationskoeffizient**

• $r_{xy} = \frac{s_{xy}}{\sqrt{s_x^2 \cdot s_y^2}}$ s'appelle **coefficient de corrélation**

6.2.4 Korrelationskoeff.: Bedeutung — Coeff. de corrélation: Signification

Untersuchung: • *Etude:*

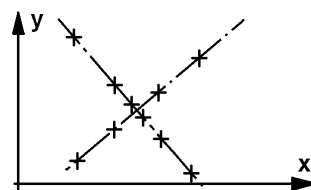
$$\leadsto r_{xy} = 0 \Leftrightarrow 0 = s_{xy} := \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x}) \cdot (y_i - \bar{y}) = \frac{1}{n-1} \sum_{i=1}^n (\Delta x_i) \cdot (\Delta y_i) =$$

$$\frac{1}{n-1} ((\sum_{i=1}^n x_i y_i) - n \bar{x} \bar{y}) \Leftrightarrow \frac{1}{n} \cdot \sum_{i=1}^n x_i y_i = \overline{\vec{x}, \vec{y}} = \bar{x} \cdot \bar{y}, \quad \vec{x} = (x_1, \dots, x_n)^T, \quad \vec{y} = (y_1, \dots, y_n)^T$$

Konsequenz: • **Conséquence:**

$$r_{xy} = 0 \Leftrightarrow \left\langle \begin{pmatrix} \Delta x_1 \\ \vdots \\ \Delta x_n \end{pmatrix}, \begin{pmatrix} \Delta y_1 \\ \vdots \\ \Delta y_n \end{pmatrix} \right\rangle = 0 \Leftrightarrow \begin{pmatrix} \Delta x_1 \\ \vdots \\ \Delta x_n \end{pmatrix} \perp \begin{pmatrix} \Delta y_1 \\ \vdots \\ \Delta y_n \end{pmatrix} \Leftrightarrow \overline{\vec{x}, \vec{y}} = \bar{x} \cdot \bar{y}$$

Diese Situation tritt auf, wenn die Punkte „wolkenartig“ verteilt sind, wie man schon mit vier Punkten sieht, die die Ecken eines achsenparallelen Rechtecks bilden. • *On a cette situation si les points sont distribués en forme de nuage, comme on voit déjà avec quatre points qui sont situés dans les sommets d'un rectangle.*



Detailliertere Untersuchung: • *Etude détaillée:*

Seien • *Soient*

$$\Delta \vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} - \bar{x} \cdot \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} - \begin{pmatrix} \bar{x} \\ \vdots \\ \bar{x} \end{pmatrix} = \begin{pmatrix} \Delta x_1 \\ \vdots \\ \Delta x_n \end{pmatrix},$$

Hier ist: • *Ici, on utilise:*

$$\varphi = \Delta(\Delta \vec{x}, \Delta \vec{y}).$$

$$\Delta \vec{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} - \bar{y} \cdot \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} - \begin{pmatrix} \bar{y} \\ \vdots \\ \bar{y} \end{pmatrix} = \begin{pmatrix} \Delta y_1 \\ \vdots \\ \Delta y_n \end{pmatrix}$$

Es gilt: • Il vaut:

$$r_{xy} = \frac{s_{xy}}{\sqrt{s_x^2 \cdot s_y^2}} = \frac{\frac{1}{n-1} \cdot \langle \Delta \vec{x}, \Delta \vec{y} \rangle}{\sqrt{\frac{1}{n-1} \cdot \langle \Delta \vec{x}, \Delta \vec{x} \rangle \cdot \frac{1}{n-1} \cdot \langle \Delta \vec{y}, \Delta \vec{y} \rangle}} = \frac{\frac{1}{n-1} \cdot |\Delta \vec{x}| \cdot |\Delta \vec{y}| \cdot \cos(\varphi)}{\frac{1}{n-1} \cdot |\Delta \vec{x}| \cdot |\Delta \vec{y}|} = \cos(\varphi)$$

$$r_{xy} = 0 \Leftrightarrow \langle \Delta \vec{x}, \Delta \vec{y} \rangle = 0 \Leftrightarrow \Delta \vec{x} \perp \Delta \vec{y} \Leftrightarrow \cos(\varphi) = 0 \Leftrightarrow \varphi = \frac{\pi}{2} + n \cdot \pi, \quad n \in \mathbb{Z}$$

$$r_{xy} = 1 \Leftrightarrow \langle \Delta \vec{x}, \Delta \vec{y} \rangle = |\Delta \vec{x}| \cdot |\Delta \vec{y}| \Leftrightarrow \Delta \vec{x} \parallel \Delta \vec{y} \Leftrightarrow |\cos(\varphi)| = 1 \Leftrightarrow \varphi = n \cdot \pi, \quad n \in \mathbb{Z}$$

Korollar: • **Corollaire:**

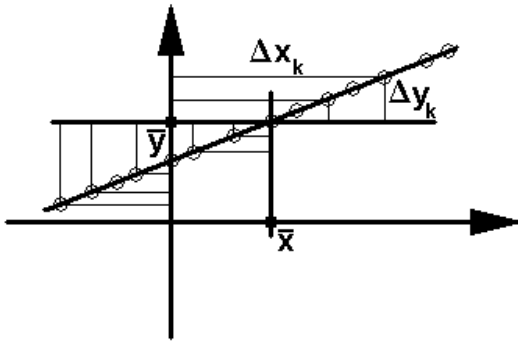
Vor.: • **Hyp.:**

$$\varphi = \angle(\Delta \vec{x}, \Delta \vec{y}), \quad \Delta \vec{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} - \bar{x} \cdot \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}, \quad \Delta \vec{y} = \begin{pmatrix} y_1 \\ \vdots \\ y_n \end{pmatrix} - \bar{y} \cdot \begin{pmatrix} 1 \\ \vdots \\ 1 \end{pmatrix}$$

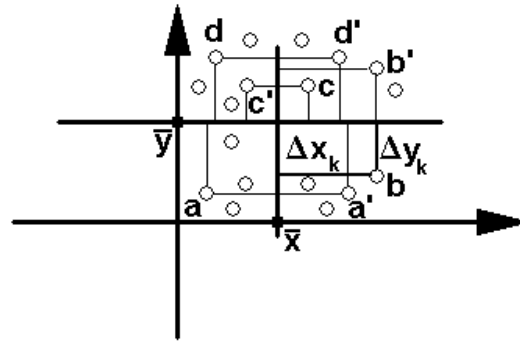
Beh.: • **Thè.:**

$$r_{xy} = 0 \Leftrightarrow \varphi = \frac{\pi}{2} + n \cdot \pi, \quad n \in \mathbb{Z}$$

$$|r_{xy}| = 1 \Leftrightarrow \varphi = n \cdot \pi, \quad n \in \mathbb{Z}$$



$$\varphi = n \cdot \pi, \quad n \in \mathbb{Z}$$



$$\varphi = \frac{\pi}{2} + n \cdot \pi, \quad n \in \mathbb{Z}$$

Bsp.: • **Exemple:**

1. $x = \{2, 3, 4, 6, 8, 12\}, \quad y = \{4, 6, 8, 12, 16, 24\} \quad (y = 2x) \Rightarrow r = 1$
2. $x = \{-2.1, -2.1, 0, 0, 2, 2.1\}, \quad y = \{4, -4, 3, -3, 4, -4\} \Rightarrow r = -0.0106422$
3. $x = \{-2.1, -2, 2, 2.1\}, \quad y = \{4, -4, 4, -4\} \Rightarrow r = -0.024383$

6.2.5 Rechnen mit Punktschätzern: Probleme — Calculer avec des estimateurs: Problèmes

Sei μ_X = Erwartungswert einer Messgrösse X , $Y := f(X)$.

• Soit μ_X = valeur attendue d'une grandeur mesurée X , $Y := f(X)$.

Sei μ_X geschätzt durch eine erwartungstreue Punkt-Schätzung $\hat{\mu}_X$.

• Soit μ_X estimé par un estimateur de vraisemblance sans biais $\hat{\mu}_X$.

$$\Rightarrow Y = \mu_Y + \Delta Y \approx f(\mu_X) + f'(\mu_X)(X - \mu_X) \text{ (Taylor).}$$

$$\leadsto \bar{y} \approx \bar{\mu}_Y = \mu_Y \approx \bar{f}(\mu_X) + \overline{f'(\mu_X)(X - \mu_X)} = f(\mu_X) + \bar{f}'(\mu_X)(\bar{X} - \bar{\mu}_X) = f(\mu_X) + f'(\mu_X) \underbrace{(\bar{X} - \mu_X)}_{=0}$$

$$\Rightarrow \bar{y} \approx \mu_Y \approx f(\mu_X)$$

Satz: • Théorème: $Y = f(X) \Rightarrow \bar{y} \approx \mu_Y \approx f(\mu_X)$

Verbesserung der Approximation: • *Amélioration de l'approximation:*

$$Y = \mu_Y + \Delta Y \approx f(\mu_X) + f'(\mu_X)(X - \mu_X) + \frac{1}{2} f''(\mu_X)(X - \mu_X)^2$$

$$\longrightarrow \mu_Y \approx f(\mu_X) + f'(\mu_X) \cdot 0 + \frac{1}{2} f''(\mu_X) \sigma_X^2 = f(\mu_X) + \frac{1}{2} f''(\mu_X) \sigma_X^2$$

Satz: • Théorème: $\bar{x} \approx \mu_X, \bar{y} \approx \mu_Y, Y = f(X) \Rightarrow \bar{y} \approx f(\mu_X) + \frac{1}{2} f''(\mu_X) \sigma_X^2$

Anwendung für Vertrauensintervalle: • *Application pour des intervalles de confiance:*

Sei • *Soit* $\mu_X \in [\bar{x} - \Delta x, \bar{x} + \Delta x], \Delta x \approx \sigma_X$ (kurz • *bref* $\mu_X = \bar{x} \pm \Delta x$).

$\leadsto$ Problem: • *Problème:* $\sigma_Y = ?$

$$Y = f(X), Y \approx f(\mu_X) + f'(\mu_X)(X - \mu_X) \Rightarrow Y - f(\mu_X) \approx f'(\mu_X)(X - \mu_X)$$

$$\Rightarrow (Y - f(\mu_X))^2 \approx (f'(\mu_X))^2 (X - \mu_X)^2 \Rightarrow \sigma_Y^2 \approx (f'(\mu_X))^2 \sigma_X^2.$$

Satz: • Théorème: $\sigma_Y^2 \approx (f'(\mu_X))^2 \sigma_X^2$

Die hier angefangenen Ausführungen sollen in den Übungen weiter verfolgt werden... Weitere Ausführungen finden sich unter dem nachstehend angegebenen Link:

• *Nous allons continuer de développer ces idées dans les exercices... On trouve des explications auxiliaires sous le lien qui suit:*

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

6.3 Zum Schluss — Quant à la fin

Damit sind die ersten Grundlagen der Statistik besprochen. In der Praxis kann man damit allerdings noch nicht viel anfangen. Von Varianzanalyse, Zeitreihenanalyse und dergleichen haben wir noch nicht gesprochen. Viele Tests sind nicht erwähnt...

• *On a ainsi discuté les premières bases des statistiques. Mais en pratique, on ne peut pas faire de grands sauts avec cela. On n'a pas encore parlé de l'analyse de variance, de l'analyse des séries chronologiques et de telles choses. Beaucoup de tests ne sont pas mentionnés...*

Problem: • Problème:

Die Normalverteilung existiert in der Realität praktisch nie, da reale Intervalle praktisch nie unendlich gross und oft auch nur positiv sind. Viele Testvariablen sind nur Approximationen. Seriöse Arbeit benötigt

ein grosses Datenmaterial... Da hilft die Fachliteratur weiter.

- *La distribution normale n'existe presque jamais dans la réalité, parce que les intervalles réels ne sont pratiquement jamais infiniment grands et souvent ils ne sont que positifs. Beaucoup de variables de test sont seulement des approximations. Pour un travail sérieux il faut un grand matériel de données... Là la littérature spécialisée va nous aider.*

Literaturverzeichnis

- [1] Fachlexikon *a b c*. Verlag Harri Deutsch, Bibliographisches Institut Mannheim, Wien, Zürich. Dudenverlag (Bibl.: A1)
- [2] Herbert Amann, Joachim Escher, Analysis I, II, Birkhäuser (Bibl.: A2)
- [3] Heinz Bauer, Wahrscheinlichkeitstheorie und Grundzüge der Masstheorie, De Gruyter (Bibl.: A3)
- [4] Bosch K., Elementare Einführung in die Wahrscheinlichkeitsrechnung, Vieweg (Bibl.: A4)
- [5] Brenner, Lesky. Mathematik für Ingenieure und Naturwissenschaftler. AULA-Verlag Wiesbaden (Bibl.: A5)
- [6] Claus, Schwill. Schüler–Duden, Die Informatik. Bibliographisches Institut Mannheim, Wien, Zürich. Dudenverlag (Bibl.: A6)
- [7] Cramer, Mathematical Methods of Statistics, Princeton (Bibl.: A7)
- [8] Iyanaga, Kawada. Encyclopedic Dictionary of Mathematics. MIT Press, Cambridge Mass., London GB (Bibl.: A8)
- [9] Erwin Kreyszig, Statistische Methoden und ihre Anwendungen, Vandenhoeck & Ruprecht Göttingen 1965/ 72 (Bibl.: A9)
- [10] Erwin Kreyszig, Statistische Methoden und ihre Anwendungen, Vandenhoeck und Ruprecht, Neuausgabe (Bibl.: A10)
- [11] Meschkowski. Mathematisches Begriffswörterbuch. BI Hochschultaschenbücher. Bibliographisches Institut Mannheim. (Bibl.: A11)
- [12] Regina Storm, Wahrscheinlichkeitsrechnung, Mathematische Statistik, Statistische Qualitätskontrolle, Fachbuchverlag Leipzig, Köln (Bibl.: A12)
- [13] L.B. van der Waerden, Mathematische Statistik, Birkhäuser (Bibl.: A13)
- [14] Mathematikkurs für Ingenieure
Teil 4 $\diamond$ Einführung in die Boolesche Algebra, Rolf Wirz, Biel 1999, (Bibl.: A14)

<http://rowicus.ch/Wir//Scripts/Teil4Bool.pdf>

- [15] Script $\diamond$ Math $\diamond$ Ing, $\diamond$ Analysis $\diamond$ Analyse $\diamond$, kurz & bündig $\diamond$ concis, Rolf Wirz, Biel 1999, (Bibl.: A15)

<http://rowicus.ch/Wir//Scripts/KAnaGdf.pdf>

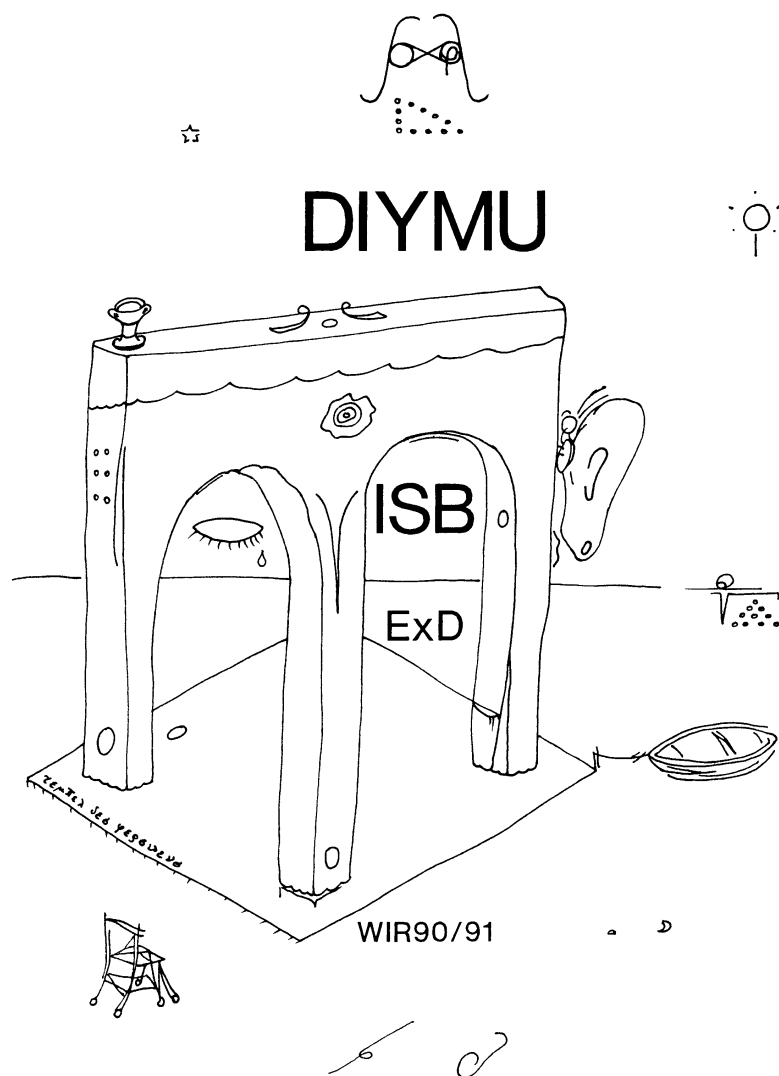
- [16] Script $\diamond$ Math $\diamond$ Ing, $\diamond$ Fortsetzung Mathematik $\diamond$ Suite mathématiques $\diamond$, kurz & bündig $\diamond$ concis, Rolf Wirz, Biel 2000/01,... (Bibl.: A16)

<http://rowicus.ch/Wir//Scripts/KursMathPubl.pdf>

- [17] Vom Autor. Mathematik für Ingenieure *Teile 1 ff* (Bibl.: A17)
- [18] Vom Autor. *DIYMU* (Do it yourself Mathematik Übungsbuch). Ingenieurschule Biel 1991 (Bibl.: A18)

Anhang A

Aus dem DIYMU



Übungen • *Excercises*

<http://rowicus.ch/Wir//TheProblems/Problems.html>

Anhang B

Fehler von statistischen Kenngrössen und Standardfehler

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

Anhang C

Monte-Carlo, Resampling und anderes

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

Anhang D

Eine Bootstrap–Anwendung Schritt für Schritt (mit *Mathematica*)

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

(Siehe auch <http://rowicus.ch/Wir/MathematicaPackages/Bootstrap.pdf>)

(Print: http://rowicus.ch/Wir/Scripts/KursWahrschStatistdf_Print.pdf)

Anhang E

Datensatzänderung

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/Zusatz/AnhangStatistDatenkorrektur.pdf>

oder

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

Anhang F

Spezielle Wahrscheinlichkeitssituationen

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

Anhang G

Hinweise zur Datenanalyse

- *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Siehe auf

<http://rowicus.ch/Wir/Scripts/KursWahrschStatistAnhangd.pdf>

Anhang H

Crashkurs Kombinatorik, Wahrscheinlichkeit: Link

• *Ici, il y a pour le moment seulement le texte allemand à disposition. Momentanément, la traduction française manque encore.*

Neuere Darstellung siehe: • *Nouvel exposé allemand voir:*

<http://rowicus.ch/Wir/Scripts/TEIL6dCrashKursWahrschKomb.pdf>

Ende *Fin*