
Poznawanie a myślenie statystyczne

1.1. Poznawanie

Jedną z najbardziej charakterystycznych cech człowieka jako gatunku *homo sapiens* jest jego zdolność myślenia. Człowiek jest istotą myślącą, *homo cogitans*, on wierzy, wątpi, ma przekonania, wnioskuje, złości się, raduje. Trudno powiedzieć, co to takiego jest myślenie, J. Bocheński określa je jako ruch duchowy (*geistige Bewegung*). Jeżeli celem myślenia jest uzyskanie wiedzy, to nazywa się ono wówczas poznawaniem. Wiedzę o świecie zewnętrznym, a także o samym sobie, człowiek pozyskuje nie tylko poprzez myślenie, ale i poprzez obserwację świata oraz przeprowadzanie eksperymentów.

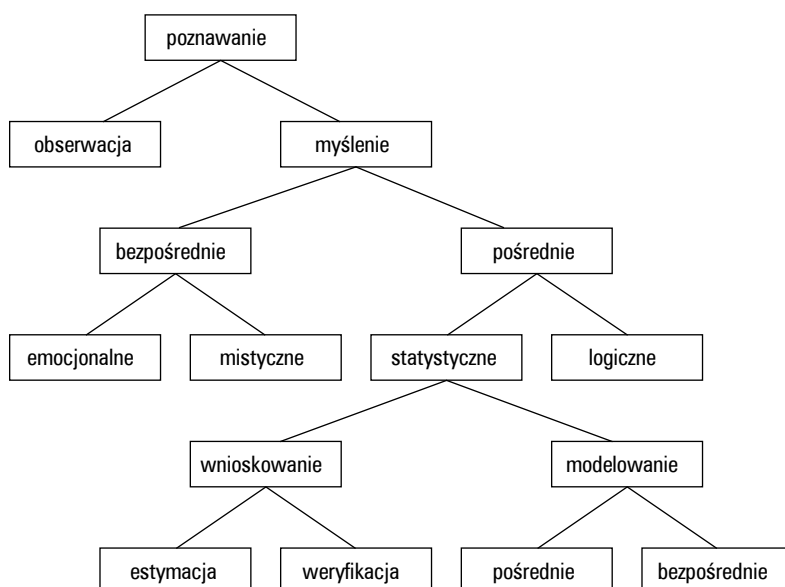
Myślenie z kolei może się odbywać z udziałem rozumu i nazywa się je wówczas rozumowaniem albo myśleniem racjonalnym, od łacińskiego słowa *ratio*, oznaczającego rozum. Rozumowanie określane jest mianem myślenia pośredniego, gdyż wykonywanie tej czynności umysłowej odbywa się za pomocą rozumu. Oprócz rozumowania istnieją też różne rodzaje myślenia bezpośredniego, stanowiące również źródło naszej wiedzy o świecie. Przykładem takiego źródła poznania jest doświadczenie mistyczne, a także przeżycia psychiczne, emocjonalne.

W dalszej części omawiane jest wyłącznie myślenie racjonalne, ale nie dlatego, że jest ono określane mianem myślenia naukowego, i z tego powodu zasługującego na szczególne traktowanie, lecz dlatego, że można się go nauczyć chociażby przez wskazanie reguł pozwalających uniknąć błędów myślenia. Istnieją przecież rzeczy, których rozumem pojąć nie można, ale można je sercem odczuć. Pascal wyraził tę prawdę w postaci słynnej sentencji: *Le cœur a ses raisons que la raison ne connaît pas* (serce ma swoje racje, których rozum nie zna).

Wyodrębnimy dwa rodzaje myślenia racjonalnego: myślenie logiczne i myślenie statystyczne. Myślenie statystyczne można traktować jako specyficzny rodzaj myślenia logicznego. Ze względu na tę specyfikę, a także ze względu

na historię jego rozwoju, myślenie statystyczne wyodrębniemy jako drugi, obok myślenia logicznego, podstawowy filar myślenia pośredniego. Wyróżnimy przy tym dwa rodzaje rozumowania statystycznego: wnioskowanie i myślenie o rzeczywistości fizycznej za pomocą modeli statystycznych, które w skrócie określimy mianem modelowania statystycznego. Modelowanie to stosowane może być w formie bezpośredniej, jako formalna imitacja zjawisk świata realnego, lub w formie pośredniej, jako narzędzie wspomagające pozastatystyczne sposoby analizy danych. Klasyfikacja sposobów poznawania przedstawiona jest na rysunku 1.1.

Rysunek 1.1. Sposoby poznawania rzeczywistości



1.2. Przedmiot myślenia statystycznego

Głównym przedmiotem myślenia statystycznego są dane statystyczne. Mają one wyraźną specyfikę odróżniającą je od wszystkich innych danych o świecie. Specyfikę danych statystycznych określają trzy następujące własności: masowość, zmienność, niepewność.

Pierwsza charakterystyka jest dość oczywista, oznacza ona bowiem, że danych jest zwykle bardzo dużo. Aby uzmysłowić sobie, co znaczy dużo, weźmy na przykład dane dotyczące średniego miesięcznego dochodu, jakim dysponują

gospodarstwa domowe w Polsce. Gospodarstw domowych jest około 13 000 000. Załóżmy, że średni miesięczny dochód jest liczbą czterocyfrową. Dane dotyczące dochodów gospodarstw domowych stanowią więc zbiór 65 milionów cyfr. Wyobraźmy sobie, jak gruba byłaby książka, gdybyśmy je zapisali, oddzielając chociażby tylko spacją jedną liczbę od drugiej.

Zmienność danych jest także dość dobrze i jednoznacznie pojmowana przez większość ludzi. Jeśli temperatura jednego dnia wynosi 5°C , drugiego zaś minus 10°C , trzeciego 0°C itd., to bez większego zastanowienia powiemy, że jest ona bardziej zmienna od takiej, gdy każdego dnia jest ok. 10°C .

Trzecia cecha danych statystycznych jest nieco bardziej skomplikowana. Pozornie jest ona również dość oczywista. Wiemy przecież, co znaczy stwierdzenie: „nie wiemy, ile jutro spadnie śniegu”. Trzeba jednak zauważyć, że istnieją bardzo różne rodzaje niepewności. Niepewność dotycząca nasilenia jutrzejszych opadów jest zupełnie inna od niepewności wygrania w totolotka, a jeszcze bardziej inna niż niepewność tego, czy wykupione dzisiaj akcje jakiejś firmy okażą się zyskowne, czy też przynoszące straty.

Niepewność dotyczy zwykle przyszłych zdarzeń. Jeżeli przyszłe zdarzenie związane jest z ewentualną stratą lub utratą czegoś, czego nie chcielibyśmy tracić, to niepewność nazywa się ryzykiem. Ryzyko jest to więc niechciane niepewne zdarzenie.

W takim przypadku niepewność ma dwa aspekty, jeden z nich dotyczy samego zajścia zdarzenia niepożądanego, drugi zaś mówi o rozmiarze, czyli wielkości straty. Niezależnie od intuicyjnego rozumienia tych trzech podstawowych charakterystyk danych statystycznych, potrzebne są sposoby ich pomiaru.

1.3. Niepewność

Specyficzność myślenia statystycznego polega głównie na tym, że dotyczy ono niepewności. Od najdawniejszych czasów człowiek dążył do absolutnej pewności w poznaniu prawdy. Dążenie takie wynikało między innymi z przekonania, że istnieją wielkie i niezmiennie prawa wytyczające ścieżki, którymi wędrujemy. To powiedzenie Goethego podzielało wielu myślicieli, których nazwiska będą się tu nieraz pojawiały. Leibniz, Poincaré, Laplace i inni uważali, że wszystko, co istnieje, ma swoją rację i przyczynę. Możliwości umysłu ludzkiego są jednak ograniczone. Niektórzy sądzą nawet, że niezależnie od postępu nauki pewnych granic umysłu nigdy nie przekroczymy. E. du Bois-Reymond określił to łacińskim powiedzeniem: *ignoramus et ignorabimus* (nie wiemy i nie będziemy wiedzieć). Nie będąc w stanie poznać i pojąć wielu przyczyn, mędrcy starożytnej Grecji wymyślili sztuczne bóstwo, nadając mu imię Τυχη (tyche), które Rzy-

mianie nazwali Fortuną. W języku polskim występuje jako Los, przy czym jest on rozumiany dwojako, jako **przypadek** i jako przeznaczenie.

Filozofowie od wielu lat toczą spór o to, czy losowość jest immanentną cechą świata, w którym żyjemy i stanowimy jego część, czy też stanowi wygodną zasłonę naszej niewiedzy o świecie. Z punktu widzenia myślenia statystycznego nie jest ważne, kto w tym sporze ma rację. Ważne jest to, że losowość ma swoje prawa, które człowiek jest w stanie poznać i wykorzystać dla swego dobra.

Gdy się okazało, że niepewności nie można się pozbyć, dołożono wiele starań, aby nauczyć się żyć zarówno w warunkach niepewności, jak i ryzyka. Zamiast zdawać się na kaprysy fortuny, zaczęto poznawać jej zachowanie i okazało się, że zachowanie to można imitować za pomocą tzw. ruletki, którą nazywa się też kołem fortuny lub loterią. Fundamentalną cechą tego niezwykle prostego i użytecznego urządzenia jest prawdopodobieństwo. Pojęcie prawdopodobieństwa stanowi fundament, na którym zbudowano gmach myślenia statystycznego. Stanowi ono nie tylko charakterystykę liczbową niepewności, ale jest też wykorzystywane do pomiaru zmienności danych statystycznych.

1.4. Probabilistyka

Słowo „prawdopodobieństwo” jest znane większości ludzi, bywa jednak rozumiane bardzo różnie, bo używa się go w różnych kontekstach. Wiadomo, że bardzo trudno jest wytyczyć granice między różnymi kontekstami. Zazwyczaj łatwo można podać przypadki typowe, dalekie od siebie. Wydzielmy więc dwie typowe sytuacje określające dwa różne rozumienia słowa „prawdopodobieństwo”. Przedstawmy je na dwóch przykładach. Załóżmy, że mamy sześciocienną kostkę do gry. Zapytajmy, jakie są szanse, że wypadnie sześć oczek, gdy tę kostkę upuścimy na podłogę?

Każdy, kto widział kostkę, odpowie, że **szanse** wypadnięcia sześciu oczek są jak 1 do 5, gdyż jest sześć ścianek, na jednej jest sześć oczek, na pozostałych pięciu ściankach jest inna liczba oczek (por. par. 1.5).

Drugi przykład dotyczy katastrofy w Smoleńsku. Zaraz po katastrofie jedni orzekli, że przyczyną katastrofy była nieodpowiedzialna decyzja lądowania na lądowisku w warunkach zupełnego zamglenia. Inni zaś uznali, że katastrofa była wynikiem misternie zaplanowanego zamachu, a mgła była sztucznie wytworzona. Przyjmijmy, że myślimy według logiki klasycznej. Prawda więc istnieje niezależnie od tego, kto i co sądzi na ten temat. Zapytajmy teraz: które z orzeczeń jest bliższe prawdy?

Ponieważ nie wiemy, jak zmierzyć odległość od prawdy, raczej powinniśmy zapytać, które z tych dwóch orzeczeń, według twego rozumu czy odczucia, jest bardziej przekonujące? Jeśli jakieś stwierdzenie jest bardziej przekonujące, to mówi się, że jest bardziej zbliżone do prawdy, czyli bardziej prawdopodobne. W tym właśnie sensie słowo „prawdopodobieństwo” różni się od słowa „szansa” wyjaśnianego w pierwszym przykładzie. Zamiast słowa „szanse”, obecnie i w tym kontekście używa się słowa prawdopodobieństwo.

W starożytnej Grecji konsekwentnie rozróżniano obie sytuacje, używając dwóch różnych słów.

Na oznaczenie tego, „co się zazwyczaj zdarza”, używano słowa *eikos*. Natomiast na oznaczenie tego, co jest dostatecznie przekonujące, używano słowa *pithanon* i jemu pochodnych. Polskim odpowiednikiem pierwszego słowa jest słowo „szanse”, zaś drugiego słowo „perswazja” lub „przekonanie”. Tłumacząc dzieła greckie na język łaciński, Cyceon użył dwóch słów: *probabilis* i *verisimile*, potraktował je jednak jako synonimy, wprowadzając tym samym pewne zamieszanie terminologiczne. Łacińskie słowo *verisimile* oznacza dosłownie „podobne do prawdy”, czyli prawdopodobne, zaś słowo *probabilis* pochodzi od słowa *aprobare*, czyli „aprobować”. Oba łacińskie słowa odnoszą się więc do sytuacji przedstawionej w drugim przykładzie, czyli do stwierdzeń zasługujących na aprobatę lub takich, które można uznać jako podobne do prawdy.

Angielskimi odpowiednikami słowa *pithanon* są słowa: *credible*, *plausible*, *reasonable*, *likely*. Zaś słowo *eikos* tłumaczone jest jako *probable*, *fair*, *equitable*. W języku niemieckim używane jest słowo określające podobieństwo do prawdy, *Wahrscheinlichkeit* (to, co się jawi jako prawda). W języku rosyjskim zaś słowo *wierojatnost'* sugeruje podobieństwo nie do prawdy, a raczej to, co zasługuje na wiarę, czyli godne wiary, ale i to, co w języku polskim określamy jako wiarygodne, w języku rosyjskim jest także związane z wiarą (*dostowiernost'*). W języku polskim mamy dwa różne słowa: w jednym podkreśla się podobieństwo do prawdy, w drugim zaś wagę tego, co jest godne wiary.

Aby odróżnić prawdopodobieństwo rozumiane jako szansa zajścia zdarzenia od prawdopodobieństwa rozumianego jako podobieństwo do prawdy orzekania o czymś, czyli wypowiadanego sądu, stosowane są odrębne pojęcia. W pierwszym przypadku używa się nazwy prawdopodobieństwo fizyczne, gdyż dotyczy zjawisk zachodzących w przyrodzie. Greckim odpowiednikiem słowa przyroda jest słowo *phisi*. Ponieważ typowym przykładem zdarzeń ocenianych za pomocą takiego prawdopodobieństwa są zdarzenia związane z rzutem kością do gry, to prawdopodobieństwo określające wynik takiego rzutu nazywane jest prawdopodobieństwem aleatorycznym lub hazardem. Nazwy takie wynikają z tego, że kostka do gry w języku łacińskim jest to *alea*, zaś w języku arabskim słowo *al zahr* też oznacza kostkę do gry.

Drugie rozumienie prawdopodobieństwa dotyczy mniemań ludzkich. Mniemanie lub opinia w języku greckim określane było jako *doksa*. Stąd też ten rodzaj prawdopodobieństwa określa się mianem prawdopodobieństwa doksatycznego. Czasem jest ono określane też ogólniej jako prawdopodobieństwo epistemiczne (od greckiego słowa *episteme*, czyli wiedza).

Pierwszą pracę na temat prawdopodobieństwa aleatorycznego napisał B. Pascal w 1654 r., opublikowano ją pośmiertnie w 1665 roku. Była to praca pt. *Traité du triangle arithmétique*, w której podany był rozkład dwumianowy, nazywany w literaturze polskiej rozkładem Bernoullego. Drugą jest napisana po łacinie w 1657 r. i często cytowana praca Ch. Huygensa. Tytuł tego 15-stronicowego tekstu, nazywanego pierwszym podręcznikiem z rachunku prawdopodobieństwa, brzmi: *De ratiociniis in ludo aleae*, co oznacza „Rachunki przy grze w kości”. Pracę tę J. Arbuthnot przetłumaczył na język angielski i była to pierwsza praca dotycząca rachunku prawdopodobieństwa opublikowana w języku angielskim.

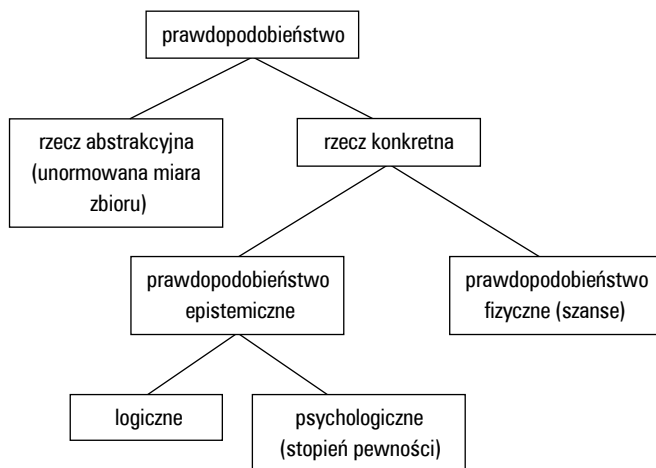
Prawdopodobieństwo fizyczne jest dość intuicyjne i dość łatwo definiowalne. W rozprawie pt. *De incerti aestimation* z 1678 r. W. Leibniz definiuje szansę zdarzenia, czyli prawdopodobieństwo w sensie aleatorycznym, w postaci ilorazu liczby przypadków sprzyjających temu zdarzeniu do liczby wszystkich możliwych przypadków. Taką definicję nazywa się dzisiaj definicją klasyczną prawdopodobieństwa i przypisuje się ją zwykle Laplace'owi.

Prawdopodobieństwo epistemiczne występuje w dwóch odmianach, jako prawdopodobieństwo logiczne i jako prawdopodobieństwo personalistyczne, nazywane też subiektywnym. W pierwszym przypadku jest ono traktowane jako miara tego, w jakim stopniu dane zawarte w przesłankach potwierdzają wniosek. Druga odmiana prawdopodobieństwa epistemicznego ma charakter behawioralny, dotyczy bowiem zachowania się osoby w warunkach niepewności. Definiowane ono jest na kilka różnych sposobów, zawsze jednak wykorzystywane jest pojęcie osoby koherentnej, czyli postępującej racjonalnie, spokojnie i mającej rozsądne preferencje, czyli takiej, która wie, czego chce i jest w tym konsekwentna.

Jeśli prawdopodobieństwo fizyczne traktować można jako właściwość zjawisk fizycznych, to prawdopodobieństwo epistemiczne traktować można jako cechę umysłu ludzkiego.

Oprócz takiego konkretnego rozumienia pojęcia prawdopodobieństwa, jest ono ujmowane zupełnie abstrakcyjnie jako pojęcie formalnej teorii matematycznej. Dlatego też nazywane jest ono prawdopodobieństwem matematycznym. Aksjomatyczną teorię takiego prawdopodobieństwa sformułował A. Kołmogorow w 1933 r. jako szczególny wariant matematycznej teorii miary.

Wyszczególnione wyżej rodzaje prawdopodobieństwa możemy więc przedstawić tak jak na rysunku 1.2.

Rysunek 1.2. Typologia prawdopodobieństw

1.5. Stochastyka

W 1662 r. ukazała się bezimiennie książka pod tytułem *La logique ou l'Art de penser*, co znaczy „Logika, czyli sztuka myślenia”. Dziś wiadomo, że autorami tego dzieła byli A. Arnauld i P. Nicole. Dzieło to powszechnie znane jest pod tytułem łacińskim jako *Ars cogitandi*, czyli sztuka myślenia.

Już w pierwszym zdaniu autorzy deklarują, że „nie ma rzeczy cenniejszej niż zdrowy rozsądek i sprawność umysłu w odróżnianiu prawdy od fałszu”. Całe dzieło poświęcone jest wyrabianiu tej sprawności umysłu i tworzeniu ogólnej teorii poznawania naukowego, czyli racjonalnego, otaczającej nas rzeczywistości. Obecnie wiemy już dobrze, że oprócz wiedzy naukowej ocenianej w kategoriach prawdy i fałszu, istnieją przekonania, opinie i domniemania oceniane w kategoriach stopnia podobieństwa do prawdy. Równie cenną rzeczą jest więc sprawność umysłu w dochodzeniu do takich mniemań, które byłyby jak najbardziej zbliżone do prawdy. Pół wieku po ukazaniu się drukiem *Ars cogitandi*, czyli sztuki myślenia, w 1713 r. opublikowana została sztuka domniemań i przypuszczania pod tytułem *Ars conjectandi*. To wybitne dzieło napisał Jakub Bernoulli (1654–1705), a zostało opublikowane osiem lat po śmierci autora przez bratanka.

Podobnie jak w „Sztuce myślenia”, tu również już na samym początku jasno i klarownie przedstawiony jest cel dzieła. Przedstawmy go, korzystając z pracy [Furstenberg]:

Jeśli czegoś jesteśmy pewni, nie mając żadnych wątpliwości, to mówimy, że to wiemy lub rozumiemy. W przeciwnym razie tylko przypuszczamy lub mniemamy. Czynić przypuszczenia o czymś oznacza mierzyć prawdopodobieństwo o tym. Dlatego też definiujemy sztukę przypuszczania, czyli stochastykę (*ars conjectandi sive stochastice*), jako sztukę pomiaru prawdopodobieństwa rzeczy tak dokładnie, jak to możliwe, która umożliwia nam we wszystkich naszych osądach i działaniach dokonywanie takiego wyboru, który jest lepszy, bardziej satysfakcjonujący, bezpieczniejszy lub pozwala na dokładniejszą analizę (por. [Furstenberg]).

Jak z tego wynika, dzieło Bernoullego można zatytułować jako *Stochastyka*, czyli *sztuka przypuszczania* a zobaczymy wtedy pełną analogię z dziełem *Logika*, czyli *sztuka myślenia*.

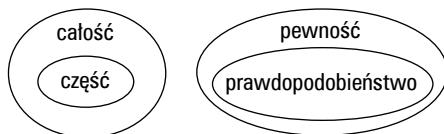
Jeśli w *Logice* jako podstawowe pojęcie wyodrębnimy pojęcie prawdy, to w *Stochastyce* rolę taką odgrywa pojęcie prawdopodobieństwa. Pojęcie to zostało zdefiniowane następująco:

Probabilitas enim est gradus certitudinis, et ab hac differt ut pars a toto

czyli prawdopodobieństwo jest to stopień pewności, od której ono się różni tak jak część od całości.

Analogia ta przedstawiona jest na rysunku 1.3.

Rysunek 1.3. Metoda porównawcza prawdopodobieństwa jako części pewności



Lata, które upłynęły między publikacjami Pascala i Bernoullego określa się jako ciemny okres prawdopodobieństwa (*dark ages of probability*).

Pojęcia stochastyka używa się obecnie prawie wyłącznie w odniesieniu do teorii procesów stochastycznych, czyli w odniesieniu do wielkości losowych traktowanych jako funkcje czasu.

1.6. Kwantyfikacja prawdopodobieństwa

Sztuka przypuszczania, którą Bernoulli nazwał stochastyką, polega głównie na określaniu prawdopodobieństwa tak, aby było ono możliwie najbliższe pewności. Słowo stochastyka pochodzi od greckiego słowa $\sigma\tau\omicron\chi\omicron\zeta$

(stochos) oznaczającego cel, odgadywanie czegoś niepewnego, zamiar lub celne określenie prawdopodobieństwa rozumianego jako stopień pewności. Przyjmijmy, że ten stopień określany będzie liczbą z przedziału $[0, 1]$ i oznaczany będzie literą p (od słowa prawdopodobieństwo), lub jako $P(E)$, gdy wyszczególnimy zdarzenie, którego prawdopodobieństwo jest oceniane.

Prawdopodobieństwo oceniane jest często w postaci **szans**. Słowo „szanse” w języku polskim pojmowane jest dwojako. Pierwsze rozumienie odpowiada angielskiemu słowu *chance* i jest wówczas traktowane jako synonim słowa „prawdopodobieństwo” (w zastosowaniu do prawdopodobieństwa fizycznego). Drugie znaczenie odpowiada angielskiemu *odds*. Jest wówczas formułowane w postaci „ułamka” postaci „ $a:b$ ” odczytywanego jako: „szanse są jak a do b ”.

Rozpocznijmy od prostego przykładu. Załóżmy, że mamy zwykłą kostkę do gry. Zapytajmy: jakie są szanse, że wypadnie sześć oczek, gdy tę kostkę upuścimy na podłogę? Kostka ma sześć ścianek, na jednej ścianie umieszczonych jest sześć oczek, na pozostałych pięciu ściankach są inne liczby oczek. Dlatego też rozsądnie jest odpowiedzieć na postawione pytanie następująco: szanse wypadnięcia sześciu oczek są jak jeden do pięciu. Często odpowiedź taką zapisuje się krócej: szanse są jak 1:5.

Na podstawie takiej odpowiedzi można uzyskać prawdopodobieństwo wypadnięcia sześciu oczek (przy jednorazowym rzucie prawidłową kostką do gry). Prawdopodobieństwo to określa się następująco:

$$\text{prawdopodobieństwo} = \frac{1}{1+5}.$$

Jeśli zdarzenie „wypadnięcia sześciu oczek” oznaczmy literą A , zaś zamiast słowa prawdopodobieństwo użyjemy jednej litery P , to powyższy „wzór”, określający prawdopodobieństwo zajścia zdarzenia A , zapiszemy krótko:

$$P(A) = \frac{1}{6}.$$

Jeśli mówimy, że szanse jakiegoś zdarzenia przypadkowego są jak $a:b$, to oznacza, że prawdopodobieństwo tego zdarzenia oceniamy jako równe ilorazowi $a/(a+b)$.

Warto w tym miejscu zwrócić uwagę na to, że w języku polskim dość często wyrażenie „**szanse są jak 1 do 6**” traktowane jest jako ułamek $1/6$. Poprawniej jest jednak użyć określenia „1 na 6”, co odpowiada też angielskiemu „1 in 6”.

Aby możliwie najprościej wyłożyć metodykę pomiaru prawdopodobieństwa, rozpatrzmy zdarzenie określające stan jutrzejszej pogody. Sformułujmy zda-

zenie: E = Jutro będzie padał deszcz. Aby określić prawdopodobieństwo tego zdarzenia, zwrócimy się do meteorologa, proponując mu dwa bilety na dwie gry hazardowe, z których musi wybrać jedną. Ponieważ tego typu gry stanowią jego główne źródło utrzymania, to możemy być pewni, że wybierze taką grę, która pozwoli uzyskać mu pewniejszy dochód.

Bilet na pierwszą grę jest następujący:

Bilet niniejszy gwarantuje jego posiadaczowi wypłatę 100 zł pojutrze, jeśli jutro będzie padał deszcz; jeśli będzie pogoda, żadnej wypłaty nie będzie.

Gra druga jest następująca:

W urnie jest 100 kul, z których 40 jest koloru białego, pozostałe są czarne.
Bilet niniejszy gwarantuje jego posiadaczowi wypłatę 100 zł, jeśli wyciągnięta przez niego kula będzie koloru białego, w przeciwnym razie żadnej wypłaty nie będzie.

Jeśli meteorolog wybierze grę pierwszą, oznacza to, że jest pewny, iż prawdopodobieństwo, że jutro będzie padał deszcz, jest większe niż 0,4. Swoje szanse uzyskania 100 zł ocenia bowiem wyżej w pierwszej grze niż w drugiej. Zakładamy oczywistą rzecz, że chce uzyskać więcej niż mniej i nie ma żadnego celu, żeby nas wprowadzić w błąd. Jeżeli chcielibyśmy ustalić bardziej dokładnie prawdopodobieństwo tego, że jutro będzie padał deszcz, to moglibyśmy proponować serię tego typu loterii, zmieniając tylko proporcje kul białych i czarnych. Cechą charakterystyczną takiego podejścia do pomiaru prawdopodobieństwa jest to, że decydent (ekspert) musi dokonać wyboru jednego z dwóch wariantów. Wyboru dokonuje na podstawie swoich **preferencji**.

Nieco inne podejście bazuje na idei stawiania (robienia) zakładów. Załóżmy, że bookmacher proponuje pewnemu graczowi zakład dotyczący jakiegoś zdarzenia niepewnego. Gracz ma określić prawdopodobieństwo p zajścia tego zdarzenia i wpłacenia kwoty pS w zamian za uzyskanie kwoty S , jeśli to zdarzenie rzeczywiście zaistnieje, i prawdopodobieństwo utraty swego wkładu w przeciwnym razie. Kwotę S ustala jednak bookmacher. Przy czym może ona być liczbą ujemną, co oznacza, że wówczas bookmacher kupuje loterię, a nie ją sprzedaje. Rozpatrzmy przykład.

Przykład 1.1**Zakład dotyczący budowy stadionu**

Założmy, że zdarzeniem niepewnym jest wybudowanie w terminie stadionu na igrzyska Euro 2012. Przyjmijmy, że gracz ocenia prawdopodobieństwo wybudowania stadionu w terminie i jest przekonany, że wynosi ono 0,6, czyli stawia 3 do 2, że stadion będzie wybudowany. Bookmacher ustala kwotę $S = 500$ zł i pobiera od gracza opłatę w wysokości 300 zł ($500 \cdot 0,6$). Jeśli stadion zostanie wybudowany, wypłaca graczowi 500 zł, jeśli nie będzie wybudowany w terminie, zatrzymuje sobie kwotę 300 zł.

Jeśliby bookmacher ustalił, że $S = -500$, to by oznaczało, że to on kupuje od gracza loterię za 300 zł w zamian za uzyskanie 500 zł, jeśli stadion zostanie wybudowany w terminie.

Rozpatrzmy jeszcze jeden przykład, który stanowić będzie podstawę do określenia osoby **koherentnej**.

Przykład 1.2**Zakład holenderski**

Założmy, że dane są cztery wykluczające się zdarzenia niepewne, łącznie zaś stanowiące zdarzenie pewne. Dla ustalenia uwagi przyjmijmy, że są to następujące zdarzenia:

E_1 = ceny akcji wzrosną o ponad 10%,

E_2 = ceny akcji wzrosną co najwyżej do 10%,

E_3 = ceny akcji się nie zmienią lub spadną nie więcej niż o 10%,

E_4 = ceny akcji spadną co najmniej o 10%.

Bookmacher prosi eksperta o oszacowanie prawdopodobieństw tych zdarzeń.

Założmy, że ekspert podaje następujące oceny:

$$P(E_1) = 0,50; P(E_2) = 0,25; P(E_3) = 0,20; P(E_4) = 0,10.$$

Bookmacher podaje kwotę do wygrania $S = 100$ zł. Ekspert wpłaca zatem łączną kwotę w wysokości 110 zł:

$$0,50 \cdot 100 + 0,25 \cdot 100 + 0,20 \cdot 100 + 0,10 \cdot 100 = 110.$$

Niezależnie od tego, które zdarzenie się zrealizuje, bookmacher wypłaca kwotę $S = 100$ zł. Czyli w każdym przypadku ekspert traci.

Założmy, że ekspert ustali następujące prawdopodobieństwa:

$$P(E_1) = 0,30; P(E_2) = 0,25; P(E_3) = 0,20; P(E_4) = 0,10.$$

Ale bookmacher, mający swobodę w ustalaniu kwot do wygrania, ustali tym razem kwotę $S = -100$ zł. Oznacza to, że to bookmacher „kupuje” loterię od eksperta.

Wyplaca mu łączną kwotę 90 zł i niezależnie od tego, które zdarzenie się zrealizuje, wygrywa kwotę 100 zł. Ekspert nie jest w tym przypadku osobą koherentną, źle szacuje prawdopodobieństwo.

Zakład tego typu, jak przedstawiony w postaci powyższego przykładu, nazywa się zakładem holenderskim. Jest to zakład, w którym jedna osoba z całą pewnością przegrywa. Jeśli osoba oceniająca prawdopodobieństwa nie godzi się na zakład tego typu, to nazywa się ją osobą koherentną. Okazuje się, że jeśli osoba będzie postępowała koherentnie, to szacowane przez nią prawdopodobieństwa będą spełniały między innymi następujące trzy warunki:

- 1) prawdopodobieństwo jest liczbą z przedziału $[0, 1]$,
- 2) prawdopodobieństwo zdarzenia pewnego jest równe 1,
- 3) prawdopodobieństwo zdarzeń wykluczających się jest równe sumie prawdopodobieństw tych zdarzeń.

Rozpatrzmy sytuację, gdy trzeba ocenić prawdopodobieństwo powtarzającego się zdarzenia losowego, na przykład oszacowania prawdopodobieństwa pogody. Załóżmy, że prawdopodobieństwo zdarzenia E wynosi p , czyli $P(E) = p$. Ocenę eksperta oznaczmy symbolem r . Określmy następującą **funkcję premijną** wyniki ekspertyzy:

$$S(r) = \begin{cases} S_1(r), & \text{gdy zdarzenie } E \text{ zajdzie.} \\ S_2(r), & \text{gdy zdarzenie } E \text{ nie zajdzie.} \end{cases} \quad (1.6.1)$$

Średnią liczbę punktów premiowych, jaką uzyska ekspert, określa się ze wzoru

$$\bar{S} = pS_1(r) + (1-p)S_2(r).$$

Ponieważ p nie jest znane, to zamiast p przyjmijmy obserwowaną częstość $\hat{p}$, z jaką realizowane jest zdarzenie E . Średnią ocenę punktową określić wówczas można w postaci wzoru:

$$\bar{S}(r) = \hat{p}S_1(r) + (1-\hat{p})S_2(r). \quad (1.6.2)$$

Jeśli $\hat{p} = r$, to mówi się, że prognozy są dobrze kalibrowane, czyli ekspert jest dobrego kalibru. Średnią ocenę określa się w takim przypadku za pomocą wzoru:

$$\bar{S}(r) = rS_1(r) + (1-r)S_2(r).$$

Przy prognozowaniu pogody stosowane są następujące funkcje:

$$\begin{aligned} S_1(r) &= 1 - (1-r)^2, \\ S_2(r) &= 1 - r^2. \end{aligned} \quad (1.6.3)$$

Średnią ocenę wyznaczymy wówczas następująco:

$$\begin{aligned}\bar{S}(r) &= \hat{p}S_1(r) + (1 - \hat{p})S_2(r) = \hat{p}(2r - r^2) + (1 - \hat{p})(1 - r^2) = \\ &= 1 - (1 - \hat{p})\hat{p} - (r - \hat{p})^2.\end{aligned}\tag{1.6.4}$$

Zauważmy, że składnik $(r - \hat{p})^2$ stanowi karę za złą kalibrację i dlatego występuje ze znakiem minus w wyrażeniu określającym średnią premię. Zauważmy ponadto, że jeśli prognozy są dobrze kalibrowane, czyli $\hat{p} = r$, to średnia premia wynosi:

$$\bar{S}(r) = 1 - (1 - r)r = 1 - (r - r^2).\tag{1.6.5}$$

Składnik $(r - r^2)$ w tym wyrażeniu także ma wyraźną interpretację, można go bowiem rozumieć jako bodziec do formułowania ocen zdecydowanych: bliskich 0 lub bliskich 1. Wyrażenie $r - r^2$ przyjmuje wartość największą dla $r = 0,5$, czyli wówczas najwięcej się odejmuje od 1 i średnia premia jest najmniejsza.

Kończąc ten rozdział, zauważmy, że kwantyfikacja, czyli pomiar, odgrywa niezwykle istotną rolę nie tylko w procesach myślenia statystycznego, ale w ogóle stanowi on podstawę wszelkiego precyzyjnego rozumowania. Powołajmy się jeszcze raz na Arbuthnota, który już w 1692 roku napisał: „Niewiele spośród znanych nam rzeczy nie da się sprowadzić do Rozumowania Matematycznego; a jeśli się tak dzieje, jest to znak, że nasza Wiedza o nich jest niewielka i nieuporządkowana (w oryginale: *There are very few things which we know; which are not capable of being reduced to a Mathematical Reasoning; and when they cannot, it's a sign our Knowledge of them is very small and confused*), a więc dużo wcześniej niż często cytowane podobne powiedzenie lorda Kelvina.

1.7. Zasady myślenia statystycznego

Podstawową zasadę klasycznej logiki stanowi tzw. zasada deterministycznej konieczności. Logiką klasyczną określa się taki system myślenia, w którym przestrzegane są następujące cztery prawa:

- 1) identity (princípio identitatis),
- 2) sprzeczności (princípio contradictionis),
- 3) wyłączonego środka (princípio exclusi medii),
- 4) racji dostatecznej (princípio rationis sufficientis).

W myśleniu statystycznym zamiast zasady deterministycznej konieczności stosowane są tzw. prawa statystyczne. Niestety, kodyfikacja i systematyzacja

wszelkich praw i zasad statystycznych nie osiągnęła jeszcze takiej formy, jaką cieszy się logika klasyczna.

Rozpatrzmy kilka zasad myślenia statystycznego, na podstawie których mogliśmy wyrobić sobie przynajmniej przybliżony sąd o specyficzności tego myślenia.

Zanim to zrobimy, wprowadźmy pewną konwencję, jaka dawniej była dość rygorystycznie przestrzegana przez statystyków w ich pracach, zaś obecnie jest nieco zaniedbywana przez młodsze pokolenia.

Według tej konwencji wielkości abstrakcyjne oznacza się za pomocą liter alfabetu greckiego. Odpowiadające im wielkości empiryczne oznaczane są za pomocą odpowiednich liter alfabetu łacińskiego.

Zwróćmy przy tej okazji uwagę na spostrzeżenie Poincarégo, że wielkości abstrakcyjne, jakimi są obiekty matematyczne, a w szczególności liczby, mogą być identyczne lub równe, spełniające relację przechodniości. Wielkości empiryczne zaś mogą być tylko nierozróżnialne i często nie spełniają warunku przechodniości. Jednym ze sposobów formułowania teorii o zjawiskach empirycznych w taki sposób, aby zachowane były podstawowe prawa matematyczne, jest wykorzystanie rozkładów prawdopodobieństwa. Oznacza to, że wyniki pomiarów (obserwacji) reprezentowane są nie w postaci liczb rzeczywistych, lecz w postaci rozkładów prawdopodobieństwa określonych na odpowiednim zbiorze tych liczb. „Liczby” tego typu zajmiemy się w rozdziale 2.7. Teraz zaś powróćmy do prezentacji głównych zasad myślenia statystycznego, dotyczących w szczególności myślenia przy użyciu takich liczb.

Jako pierwszą zasadę myślenia statystycznego przyjmijmy zasadę głoszącą, że każde zjawisko jest swoistą mieszaniną dwóch skrajności: determinizmu i indeterminizmu. Wszystko to, co obserwujemy, jest złączeniem dwóch bytów idealnych, czyli abstrakcyjnych:

$$Y \approx \mu + \xi. \quad (1.7.1)$$

Wielkość stała (stabilna) μ w języku matematycznym może być określona w postaci liczby, macierzy, funkcji, zbioru lub innego obiektu formalnego. Wielkość niestabilną i niepewną ξ nazwiemy wielkością losową. Potraktujmy ją na razie jako pojęcie pierwotne.

Jeśli przyjmiemy, że jednym z naszych celów jest poznawanie świata, w którym żyjemy i działamy, to w tym przypadku możemy sformułować bardziej konkretne zadanie, a mianowicie wykrywanie regularności i poznawanie praw przyrody na podstawie tego, co empirycznie obserwujemy. Dlatego też, mając takie obserwacje $y_1, y_2, \dots, y_n, \dots$, chcemy „wyłuskać” z nich, czyli wyodrębnić, część regularną. Aby to było możliwe, trzeba przyjąć kolejną zasadę metodologiczną umożliwiającą oczyszczanie danych z tego, co nazywamy zakłóceniami losowymi.

Przyjmijmy więc, że część stabilna i regularna zajmuje zawsze miejsce centralne wśród wszystkich obserwacji.

Wielkość centralna jest najważniejsza, czyli ma najwyższą wagę wśród wszystkich obserwacji. Przyjmijmy, że to, co ma wyższą wagę, jawi się nam częściej od tego, co jest mniej ważne, mniej typowe. To, co się zwykle zdarza, Arystoteles nazwał prawdopodobnym. To, co się częściej zdarza, jest więc bardziej prawdopodobne od tego, co się zdarza rzadziej. Określmy „różnicę” między regularnym i nieregularnym w sposób następujący:

$$\rho = \mu - \xi.$$

Wagę tej różnicy, czyli wielkości odchylenia od części regularnej, określmy w postaci wzoru:

$$\beta = e^{-\pi\rho^2}. \quad (1.7.2)$$

Potraktujemy go jako drugą zasadę podstawowego trzonu myślenia statystycznego.

Zauważmy, że zasadę tę odkrył A. de Moivre, traktując ją jako zagadnienie teologiczne. Określił on bowiem w taki sposób częstość odchyłeń od pierwotnego planu ustalonego przez Stwórcę (w oryginale: „*For him it was a theological problem, he was determining the frequency of irregularities from the Original Design of the Deity*” – w ten sposób wyraził to K. Pearson).

Jeżeli chcemy praktycznie wyznaczyć stabilną wielkość centralną na podstawie obserwacji empirycznych, to musimy umieć mierzyć odchylenia tej wielkości od wszystkich pozostałych. Przyjmijmy, że $d(y_i, c)$ oznacza odchylenie (lub oddalenie) pewnej wielkości c od wielkości y_i , zaś D oznacza sumę wszystkich odchyłeń, to wielkością centralną nazwiemy wielkość m spełniającą warunek:

$$D(y_i, m) \leq D(y_i, c) \text{ dla dowolnego } c. \quad (1.7.3)$$

Często, gdy obserwacje są liczbowe, miarę odchylenia określa się następująco:

$$D(y_i, c) = \sum_i |y_i - c|^2.$$

Wielkością centralną wśród n obserwacji jest wówczas wielkość:

$$m = (y_1, y_2, \dots, y_n) / n.$$

Trawestując Einsteina, powiemy, że natura jest skomplikowana, ale złożliwości nam nie wykazuje, możemy więc ją poznawać, jeśli nawet tylko do pewnego stopnia. Jeśli tylko mamy dostatecznie dużo obserwacji, to na ich podstawie

możemy wykrywać prawidłowości, jakie rzeczywiście w naturze istnieją. Innymi słowy, możemy wyznaczać empiryczne odpowiedniki pojęć teoretycznych.

Jako trzecią zasadę przyjmijmy zatem następującą relację:

$$m \approx \mu. \quad (1.7.4)$$

Tę zasadę z pewnością można przypisać Bernoullemu. On to bowiem, według jego własnych słów, nad jej odkryciem pracował przez 20 lat i nazwał ją złotym prawem.

Idee Newtona o wszechobecnej boskości w przyrodzie gwarantującej stabilność wartości średnich można w tym przypadku traktować jako podstawę filozoficzną tej zasady. Według słów Pearsona: *Newton's idea of an omnipresent activating deity, who maintains mean statistical values, forms the foundation of statistical development.*

Kończąc ten rozdział, zauważmy, że myślenie statystyczne jest o wiele bardziej różnorodne i ma wiele innych aspektów niż te, dla których podstawę stanowią podane tu zasady.

Te trzy zasady określają główny trzon myślenia statystycznego, od niego wywodzą się inne, ważne gałęzie nauki statystycznej. Jedną z nich stanowi teoria wartości ekstremalnych.

Wartość średnią M. Cauchy określił jako tę, która znajduje się między dwiema wartościami ekstremalnymi: najmniejszą i największą. Okazuje się, że tak samo jak wartość środkową, wartości ekstremalne można również identyfikować na podstawie obserwacji empirycznych. Zasada druga jest wówczas nieco inaczej interpretowana, a mianowicie rozpatrywane są odchylenia nie od wartości centralnej, lecz od wartości ekstremalnej.

Mimo krytyki ze strony przedstawicieli statystyki odpornej, zasada ta jest przez nich także stosowana. Przytaczany przez nich aforyzm, że wszyscy wierzą w to, że błędy mają rozkład normalny (praktycy sądzą, że to twierdzenie matematyczne, a matematycy, że jest to fakt ustalony eksperymentalnie, zaś przez P. Hubera nazwany dogmatem normalności), nie stanowi żadnej podstawy do negowania zasady rozkładu normalnego. Ma ona zastosowanie nie tylko do rozkładu błędów. Już Edgeworth zwracał uwagę na istotną różnicę między pojęciami błędu i odchylenia. Załóżmy na przykład, że pewna osoba waży dokładnie 65 kg. Dokonujemy czterech pomiarów i uzyskujemy następujące wyniki: 64,1; 64,7; 65,5; 65,7. Średnia arytmetyczna wynosi 65. Załóżmy teraz, że dokonujemy dokładnych pomiarów czterech różnych osób i uzyskujemy te same wyniki. W tym przypadku wartość „centralna” wynosi także 65, chociaż takiej osoby w ogóle nie ma. Odchylenia od pewnej hipotetycznej wielkości mogą mieć również rozkład normalny. Dzięki takiemu spostrzeżeniu A. Quetelet odkrył wiele ciekawych praw biologicznych i społecznych.