

PRAVDĚPODOBNOST A STATISTIKA NA STŘEDNÍ ŠKOLE



matfyzpress

vydavatelství Matematicko-fyzikální fakulty
Univerzity Karlovy v Praze

Všechna práva vyhrazena. Tato publikace ani žádná její část nesmí být reprodukována nebo šířena v žádné formě, elektronické nebo mechanické, včetně fotokopí, bez písemného souhlasu vydavatele.

© (eds.) Jaromír Antoch, Danil Hlubinka a Ivan Saxl, 2005

© MATFYZPRESS, vydavatelství Matematicko-fyzikální fakulty
Univerzity Karlovy v Praze, 2005

ISBN 80-86732-23-1

Předmluva

Tato publikace je s výjimkou úvodní stati o statistickém myšlení sbírkou textů přednesených v akademickém roce 2003/2004 na didaktickém semináři, který pořádá již čtvrtým rokem katedra didaktiky matematiky MFF UK pro profesory středních škol (gymnázií a středních odborných škol). Vedle vyučujících matematice převážně z Prahy a Středočeského kraje se jej účastní také doktorandi, někdy studenti, a samozřejmě přicházejí i kolegové z MFF. Cílem semináře je poskytnout jeho účastníkům určitý nadhled nad tematickými celky zařazenými do středoškolské matematiky.

V akademickém roce 2003/2004 se katedra didaktické matematiky starala o organizační záležitosti semináře, zatímco jeho obsah plně zajistila katedra pravděpodobnosti a matematické statistiky. Bylo uspořádáno celkem osm seminářů, jejichž názvy jsou shodné s názvy zde publikovaných statí. Cyklus seminářů *Pravděpodobnost a statistika na střední škole* získal akreditaci Ministerstva školství, mládeže a tělovýchovy ČR číslo 29528/2003-25-262, a jeho účastníci obdrželi osvědčení o absolvování kurzu.

V posledních letech se v pedagogice věnuje pozornost představě, že každý vědní obor, resp. každý obor lidské činnosti, od nás vyžaduje jistou specifickou formu myšlení, usnadňující rychlé dosažení kvalitních výsledků. Úkolem základní a střední školy je pak spíše než hromadění konkrétních znalostí výchova k základním formám myšlení o přírodních i společensko-vědních jevech. Internetový vyhledávač Google může poskytnout kvantitativní představu o pozornosti, která je obecně jednotlivým formám myšlení věnována. 3. 5. 2004 jsme mu zadali základní hesla „typ myšlení“ a dostali následující počty odkazů uvedené vždy v závorce: umělecké (4), chemické (274), fyzikální (916), biologické (966), geometrické (2090), pravděpodobnostní a stochastické (2993), hudební (4260), historické (13 000), statistické (16 100), politické (29 200), matematické (48 300) a nejvyšší počet získalo myšlení logické (84 400), jež však je mezioborové a jen málo úspěšně mu konkuruje údajně nesprávně opomíjené myšlení chaotické či intuitivní (570). Výsledky této „soutěže“ naznačují, že obor didaktického semináře pro letošní školní rok nebyl vybrán špatně.

Složitou problematiku statistického myšlení naznačuje úvodní stať, která nebyla přednášena a je určena k obecnému zamyšlení. Historicky se pravděpodobnost rodí společně s kombinatorikou aplikovanou na lingvistiku, společenské jevy, otázky stvoření světa a posléze též na hry (i když jen těmi málokterými, kteří nedávali přednost šťastné *Náhodě*).

Další příspěvek je věnován tvorbě kalkulu diskrétní pravděpodobnosti a objasnění vztahu mezi náhodným experimentem a jeho matematickým modelem. Pojmy (podmíněná) pravděpodobnost, náhodné veličiny a jejich závislost či nezávislost, jejich střední hodnoty a rozptyly, vyrůstají z příkladů klasické pravděpodobnosti. Text *vrcholí* výkladem smyslu zákona velkých čísel a geneze Gaussovy křivky v modelu binomického rozdělení pravděpodobností. Čtenář by měl ovládat elementární kombinatoriku, měl by být ochoten chápat matematiku nejen jako *zajímavou hru* s pojmy, ale i jako nástroj pro modelování náhody přítomné v našem světě.

Jestliže se oprávněně o oboru pravděpodobnosti a matematické statistiky říká, že je *Popelkou matematiky*, pak geometrická pravděpodobnost je Popelčinou Popelkou. Je tomu tak přesto, že pravděpodobnostními jevy geometrické povahy jsme doslova obklopeni na každém kroku. Stať věnovaná této problematice se skládá ze dvou částí: první obsahuje vesměs historické úlohy vhodné pro rozšíření středoškolské výuky, druhá pak tyto úlohy zobecňuje a uvádí některé velmi obecné poznatky související s reprezentací náhodných geometrických objektů a jejich souborů prostředky geometrického výběru.

Teorie pravděpodobnosti je základem pro modelování náhodných procesů, vhodných pro popis dynamických dějů kolem nás. Náhodné procesy jsou systémy náhodných veličin realizujících se v diskrétním nebo spojitém čase a s různě silnou vzájemnou vazbou. Z mnoha tříd náhodných procesů jsou v páté přednášce vybrány Markovovy řetězce, které se přes svou poměrnou jednoduchost hodí na řadu reálných situací. Je pojednáno o základní klasifikaci stavů, čtenář se poučí, jak simulovat realizaci řetězce na počítači. Významnou limitní vlastností některých Markovových řetězců je existence stacionárního rozdělení; zde je krátce vyšetřována rychlost konvergence, klíčová pro užití metod markovského Monte Carlo.

Další stať je úvodem do klasického statistického myšlení. Ukazuje, jak porovnat hypotézu s realizovaným experimentem, jak srovnat dva náhodné výběry, a co dělat, chceme-li v předpokládaném modelu odhadnout neznámý parametr předpokládaného modelu. Hlavní část příspěvku je věnována problematice testu statistické hypotézy včetně vysvětlení základních pojmů jako jsou hladina, síla či kritický obor testu. Právě pochopení této statistické metodologie je klíčem k úspěšnému zvládnutí statistického myšlení.

Snad nejpoužívanější statistický model je regresní analýza umožňující modelovat závislost jedné náhodné veličiny (přesněji její střední hodnoty) na jedné či několika veličinách jiných. Předpokládán je model, podle něž je tato závislost lineární v neznámých parametrech. Na příkladu se školskou problematikou se ukazuje způsob odhadu neznámých parametrů i testování hypotéz o nich. Takto lze dokazovat existenci závislosti, predikovat budoucí hodnoty

vysvětlované veličiny či porovnávat způsob závislosti v několika souborech. Je naznačena regresní diagnostika, která umožňuje zhodnotit vhodnost použitého modelu k reálným datům.

Pojistná matematika je zajímavá matematická disciplína. Její teoretická část se označuje jako teorie rizika a využívá pokročilých postupů teorie pravděpodobnosti a matematické statistiky. Její praktická část je matematika denní praxe v pojišťovnách, zajišťovnách, penzijních fondech a jiných pojišťovacích institucích. Má podporu v pojistné legislativě (např. institut odpovědného pojistného matematika) a její znalost je zárukou zajímavého uplatnění v praxi.

Poslední příspěvek je věnován využití moderního CAS systému *MuPAD* ve výuce matematiky na střední škole. Připomeňme, že CAS je zkratka anglického *Computer Algebra Systems*, což v překladu znamená *systemy pro počítačovou algebru*. Již z názvu je patrné, že se jedná o programy umožňující provádět symbolické matematické výpočty na počítači. Pozornost byla soustředěna na program sice méně známý, než-li je Mathematica či Maple, přitom však velmi kvalitní, *MuPAD Light 2.0*. Tento program je možno získat *zdarma* pro vzdělávací účely.

Autoři sborníku doufají, že jeho vydání přispěje ke zvýšení zájmu o náhodné jevy a jejich popis i k přiblížení dosud vzdáleného cíle, aby podíl pravděpodobnosti a matematické statistiky ve vzdělávacím procesu byl srovnatelný s jejich významem pro postižení přírodních a společenských jevů.

Na závěr bychom chtěli poděkovat nejenom MŠMT ale též výzkumnému záměru MSM 0021620839, bez jehož finanční podpory by tento sborník nevyšel.

Jaromír Antoch, Daniel Hlubinka a Ivan Saxl

Obsah

Ivan Saxl	
<i>Statistické myšlení a jeho výuka</i>	1
Josef Štěpán	
<i>Základy pravděpodobnosti</i>	17
Ivan Saxl	
<i>Geometrická pravděpodobnost</i>	47
Viktor Beneš	
<i>Náhodné procesy</i>	73
Daniel Hlubinka	
<i>Základy statistického myšlení</i>	85
Karel Zvára	
<i>Regrese</i>	109
Tomáš Cipra	
<i>Demografie a pojistná matematika</i>	127
Jaromír Antoch, Michal Čihák a Jan Prachař	
<i>Použití programu MuPAD ve středoškolské výuce</i>	147

STATISTICKÉ MYŠLENÍ A JEHO VÝUKA

Ivan Saxl

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

ivan.saxl@mff.cuni.cz

Občan žijící v moderní „informační“ civilizaci je vystaven neustálému tlaku nejrůznějších dat i jejich interpretací a je od něj požadováno, aby o ně opíral svá každodenní rozhodnutí. Uvedme alespoň nejběžnější příklady těchto informací: veřejné mínění o různých problémech počínaje zavedením trestu smrti a konče doporučeními k nákupům, politické preference, údaje o zločinnosti a dopravních nehodách, podmínky a přednosti různých uložení úspor, data o nemocnosti, populačních tendencích, nezaměstnanosti, daňovém zatížení, pracovních příležitostech, vzdělávacím procesu a společenské objednávce profesí. Řada běžně se vyskytujících situací, k nimž se tato data vztahují, demonstruje malou připravenost občanů ke kvalifikovanému zacházení s tímto informačním tlakem (neboť mnoho informací občany o něčem přesvědčuje a k něčemu je nabádá): všichni klienti zkrachovaných peněžních ústavů, většina účastníků dopravních kolizí i mnozí nezaměstnaní a nemocní se rekrutují z těch, kteří dostupná data buď přijímali nekriticky nebo naopak je paušálně odmítli resp. se o ně vůbec nezajímali. Tato společenská situace je impulsem k uvažování o efektivitě našeho vzdělávacího procesu ve všech oblastech, pochopitelně včetně matematiky a především pravděpodobnosti a statistiky. Náš každodenní život je totiž nepřetržitým řetězem rozhodování v nejistých situacích na základě neúplných dat, zatím co rodinná i školní výchova jsou přísně kauzální, tj. pevně svazující jevy s jejich příčinami, tedy i naše jednání s jeho důsledky („Nesahej na kamna, spálíš se!“ atd.).

V posledním desetiletí minulého století dochází přímo k erupci výzkumu a prací, týkajících se fenoménu nazývaného *statistickým myšlením* a způsobů jeho kultivace jak pro běžné občany, tak pro profesionální statistiky. I když statistické vzdělávání ve školách v různé míře a orientaci probíhá již více než tři století, o jeho podstatě a složkách nebylo dosud nikdy dostatečně podrobně uvažováno. Proto není nijak překvapivé, že mezi vědci a pedagogy zabývajícími se touto problematikou neexistuje žádná názorová jednota a současný

Klíčová slova: Statistické myšlení, výuka statistiky, intuitivní přístup, práce s žáky

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839 a grantu GAČR č. 201/03/0946.

stav lze charakterizovat jako silně nesourodou tříšť více či méně správných postřehů bez hierarchizace poznatků a bez jednoznačných doporučení pro výuku. Je třeba zdůraznit, že pojem statistického myšlení zahrnuje jak pravděpodobnostní úvahy, tak interpretaci statistických dat a termín *statistické* je pouze zkratka.

V této úvodní stati postgraduálního kurzu statistiky bude nejprve stručně shrnuta historie statistické výuky v uplynulých stoletích a poté podán přehled současných představ o podstatě statistického myšlení i o psychických zábranách (snaha o kauzální interpretaci dat a ovlivnění vlastními emocemi a přáními) při jeho vytváření a uplatňování v běžném životě. Odtud pak vyplynou doporučení pro výuku statistice na základních a středních školách.

1. Výuka statistiky v minulých stoletích

Slovo statistika je italského původu, odvozeno od latinského slova *status* s významem *stav*; jeho italský tvar je *stato*. Jako první je patrně použil Girolamo Ghilini (1589 - 1669) v nevydané práci *Ristretto della civile, politica, statistica e militare scienza*. Označuje zde původně nenumерické znalosti o státu, jeho obyvatelstvu, životních podmínkách, právu, obchodu a výrobě. Až v polovině XVIII. stol. při snaze jednotlivé státy vzájemně porovnat se objevují číselné údaje a posléze se stávají hlavním obsahem tzv. *tabulkové statistiky* rozvíjející se intenzivně s nástupem absolutistických vlád a v souvislosti se dvěma idejemi rozvoje státu, jež si obě vynucovaly shromažďování dat. Síla a perspektiva státu byla odvozována z počtu obyvatel a jeho růstu. *Populacionisté* se proto dožadovali státní podpory obyvatelstva včetně jeho chudších vrstev a byli přesvědčeni, že zvýšený počet obyvatel automaticky zajistí růst výroby a obchodu. Naproti tomu *kameralisté* preferovali podporu výroby a obchodu, předpokládající, že jejich rozvoj a z něj plynoucí vyšší životní úroveň automaticky povedou k růstu počtu obyvatel. Obsahem „statistiky“ se stávají vedle počtu obyvatel také údaje o náboženství, školství, právu, armádě a průmyslu.

Výuka statistiky začíná v roce 1660, kdy Hermann Conring začíná na universitě v Helmstadtu přednášet vědu o státě (*Staatenkunde*). „Otcem statistiky“ je nazýván G. Achenwall, profesor nejprve v Marburgu (1746) a poté v Göttingen. Achenwall již používá termín statistika, který podle vlastního tvrzení odvozuje od italského *statista*, znamenajícího *státník*. Z Německa se výuka statistiky šíří v šedesátých letech XVIII. stol. do sousedního Rakouska (1769) a po napoleonských válkách v souvislosti s jeho územními zisky i do Itálie (Pavie 1814, Padova 1815). V roce 1807 začíná výuka statistiky v Nizozemí a o rok později také v Belgii.

Zpočátku přísně nenumerická statistika, zvaná *univerzitní* pro rozlišení od úřední *tabulkové*, je součástí právního a posléze i historického vzdělání, často je citován výrok A. L. von Schölzera, Achenwallova nástupce v Göttingen: *Geschichte ist eine vorlaufende Statistik; und Statistik ist eine stillstehende Geschichte* (Dějiny jsou nepřetržitě se měnící statistikou a statistika jsou zastavené dějiny). Finanční obtíže monarchií si ovšem vyžadují nové peněžní zdroje a význam úřední tabulkové statistiky roste. Rakouské školy se stávají statistickými centry, profesori statistiky a politických věd s pomocí místních úřadů shromažďují provinční data, na jejichž základě se vytváří statistika celého státu. Podobně je tomu v celé Evropě; v r. 1800 vzniká v Paříži Bureau de la Statistique Generale, Pruský statistický úřad začíná svou činnost v r. 1805. Velkou zásluhu o rozvoj statistiky má belgický astronom L. A. J. Quetelet (1796 - 1874), který začal aplikovat numerické postupy na různé stránky společenského života i na jednotlivé občany. Dodnes se používá Queteletův index obezity QI (též BMI z anglického **b**ody **m**ass **i**ndex, rovný váze v kilogramech dělené čtvercem výšky v metrech); $QI > 30$ popisuje „úřední obezitu“. Jeho činnost vyvolala řadu souhlasných i nesouhlasných reakcí, nicméně vedla k zakládání nových statistických organizací (Londýnská statistická společnost 1834, Americká statistická společnost 1839, Rakouský statistický úřad 1840, Ruská zeměpisná společnost 1849) a jsou pořádány Mezinárodní statistické kongresy (1851 až 1876), marně se pokoušející o všeobecné zavedení metrického systému.

Ke změnám dochází i na školách. V Belgii dochází v r. 1835 k přechodu statistické výuky na fakultu filosofickou, v r. 1849 je však úplně zrušena a zůstávají pouze přednášky z pravděpodobnosti, probíhající od r. 1835 na fakultě přírodovědecké. Na fakultě filosofické se však statistika spolu se zeměpisem vyučuje v Savojsku, odkud se po sjednocení Itálie šíří po celé zemi. Ve Spojených státech zahajuje výuku v r. 1845 virginská universita v Charlottesville v rámci morální filosofie (etiky), v roce 1873 je následována Yaleskou universitou (nový předmět se nazýval „Veřejné finance a průmyslová statistika“) a do r. 1900 se statistika objevuje na řadě dalších universit, včetně MIT, Harvardu aj. Obsahem byla státní ekonomie, občanské záležitosti a právo (již Quetelet věnoval velkou pozornost statistice zločinnosti a stupni ovlivnění jedince jeho postavením ve společnosti a gramotností), demografie aj. Charakteristické znaky náhodných jevů, především jejich variabilita, byly zanedbávány a pozornost věnována především popisovaným jevům a hledání analogií k fyzikálním zákonům, v nichž by střední hodnoty, odvozené z dat a relativního zastoupení jejich různých složek, hrály roli fyzikálních konstant. To je ostatně dodnes přetrvávající laická představa o statistice.

Teprve na sklonku století přicházejí s novými výukovými koncepcemi F. Galton a K. Pearson v souvislosti s aplikacemi statistiky v biologii. Pod jejich vlivem se rozvíjí biometrická statistická škola, považující biologické charakteristiky za nezávislé stejně rozdělené jevy a věnující pozornost pravděpodobnostním rozdělením, variabilitě jevů, regresi a statistickým testům. K ní patří také G. U. Yule, který ovlivnil několik generací statistiků svým opakovaně vydávaným *Úvodem do statistické teorie* a zasloužil se mj. o rozvoj statistiky literárních textů. Ve dvacátých a třicátých letech 20. století vytváří R. A. Fisher statistiku malých výběrových souborů s aplikacemi v zemědělství. Zároveň je znovu navázáno úzké spojení s teorií pravděpodobnosti (termín *matematická statistika* se objevuje v sedmdesátých letech XIX. stol. v Německu), jež bylo koncem XVIII. století přerušeno na základě představy, že teorie založená na hrách nemůže být aplikována na společenské jevy. Díky axiomatickému základu vybudovanému Kolmogorovem ve třicátých letech XX. století, a patrně též zásluhou matematické formulace četnostní teorie předložené R. M. E. von Misesem, se pravděpodobnost se statistikou stávají obecně uznávanou a respektovanou částí matematiky. Od té doby je statistika pevnou součástí vysokoškolského vzdělání (přírodovědecké fakulty, sociální, ekonomické i technické školy), velmi zvolna se prosazuje i v základním školství. Stav a rozsah výuky je ovšem velmi různý a zhusta se jedná pouze o návody na formální zpracování datových souborů a provedení testů nejjednodušších hypotéz. Na základních školách bývá statistika často tou částí osnov, kterou se z časových důvodů nepodaří podrobněji probrat, mnozí učitelé ji neradi učí a její obliba u žáků je minimální. Ani vztah společnosti ke statistice není kladný, všeobecně je rozšířen názor, že je prostředkem k důkazu libovolného tvrzení, a skutečnost, že všechny statistické výroky mají pouze pravděpodobnostní charakter, není obecně známá. Naopak jsou oblíbená kategorická tvrzení „podpořená“ statistickými daty bez jakékoliv statistické interpretace a často i bez formálních testů, jen na základě pohledu do tabulky či na obrázek. Smutné důsledky této situace byly již zmíněny v úvodu, nejsou však pouze naším problémem a nelze je vyřešit rychle.

Instruktivní je přehled aktivit v oblasti výuky statistiky ve Spojených státech amerických v XX. století (Burril, 1991). O začlenění této výuky do osnov matematiky usilovali pedagogové již ve dvacátých letech; přitom požadovali, aby absolventi rozuměli aspoň jednoduchým statistickým zpracováním dat a byli schopni je interpretovat. Požadavek byl opakován ve čtyřicátých i padesátých letech, ale teprve v r. 1967 byla vytvořena komise složená ze zástupců učitelů a statistiků. Výsledkem její práce bylo několik publikací a učebnic pro žáky i učitele, pořádání kurzů pro učitele a vypracování tříletého *Programu pro zlepšení gramotnosti ve školách*. Přes zlepšení celkového

stavu, v roce 1989 pouze jediný stát (Wisconsin) vyžadoval u všech vysokoškolských studentů elementární statistickou gramotnost (v rozsahu odpovídajícím jednomu semestru). V současné době pracuje Internetové vzdělávací centrum K12 (<http://www.enc.com>) pro výuku matematiky ve 21. století, poskytující informace o literatuře, aktivitách, osnovách atd., kde lze získat představu o rozsahu vyučování pravděpodobnosti a statistiky v USA.

Podrobnou studii historie statistického vzdělávání v našich zemích do r. 1848 je kniha Závodského (1992). Stručným přehledem historie statistické výuky je práce Ottaviani (1991).

2. Zásady statistického myšlení

Úzké spojení mezi teorií pravděpodobnosti a matematickou statistikou, a především jejich pevné začlenění do matematiky, je právě jedna ze skutečností, jejíž výhodnost a účelnost pro obě disciplíny zkoumající náhodné jevy je v poslední době zpochybňována; sám termín statistické myšlení (v jeho nezkráceném smyslu, tj. pravděpodobnostní a matematicko-statistické) je produktem těchto pochyb. Matematika se svou axiomatickou stavbou a z ní vyplývajícími bezesporně dokazatelnými tvrzeními je totiž vědou deduktivní a striktně deduktivní přístup není patrně ideálním východiskem pro porozumění společenským, biologickým a patrně i mnohým fyzikálním jevům, o nichž teorie pravděpodobnosti a matematická statistika na základě velmi rozmanitými postupy získaných dat něco svou charakteristickou řečí vypovídá. Ze směsice názorů, nápadů a pozorování jsou v následující části uvedeny ty, které mají obecnější platnost, i když jsou odvozeny podrobnějším zkoumáním v omezených oblastech jevů, především v oblasti výroby a řízení. Z důvodů stručnosti nejsou uváděny doslovné citáty, pouze autoři příslušných myšlenek. Často užívaný anglický termín *variation* je v dalším překládán slovem *variabilita*, přičemž se jím míní proměnlivost jevů při zhruba stejném průměrném charakteru (např. denní produkce, změny počasí apod.).

Statistické myšlení je považováno za základní myšlenkovou aktivitu zpracovávající informaci v nedeterministickém prostředí. To konkrétně znamená, že ne každý jev má svou příčinu a že některé jevy lze (či je nutné?) považovat za působení náhody. Variabilita není nový koncept. Nové je uvědomění její existence, které se promítá do každé složitější lidské činnosti a ovlivňuje všechny aspekty řízení společnosti. Statistické myšlení vychází z uvědomění variability dějů jako základního vždy přítomného prvku, jemuž má být zpracování dat podřízeno v tom smyslu, aby se co nejvýrazněji projevil. Po jeho kvantifikaci by mělo následovat vysvětlení (G. Moore, 1990).

Podle Boxe, Huntera a Huntera (1978) je statistické myšlení zpětnovažební smyčkou následujícího typu: vyčerpávajícím způsobem jsou zpracována data, na základě indukce je předložena hypotéza, ta je deduktivně ověřena a jsou formulovány důsledky, jež jsou konfrontovány s výchozí hypotézou. Není-li plně potvrzena, je pozměněna a cyklus se opakuje. Přitom je důležité, aby indukci předcházelo podrobné seznámení se situací, kterou má úloha vyřešit, včetně průvodních okolností a způsobu získání dat.

Snee (1990) formuluje podstatu statického myšlení jako myšlenkový proces, uvědomující si přítomnost variability v každé naší činnosti, a dále skutečnost, že každá činnost je výsledkem posloupnosti vzájemně souvisejících procesů. Při konkrétní aplikaci na řízení společenských činností pak konstatuje, že jejich identifikace, charakterizace, kvantifikace, kontrola a omezení variability vedou ke zlepšení výsledku (konkrétně např. ke snížení výrobních nákladů či k levnějšímu operativnímu řízení).

Na základě této definice Americká společnost pro kvalitu uvádí ve svém slovníku statistických termínů (Glossary, 1996), že učení (statistickému myšlení) a činnost jím řízená jsou založeny na následujících principech:

- každá činnost probíhá jako systém vzájemně propojených procesů,
- variabilita se vyskytuje u všech procesů,
- její pochopení a redukce je klíčem k úspěchu.

Mallows (1998) upozorňuje na výchozí bod každé statistické analýzy, jímž je zjištění, co je v datech podstatné, jak se vztahují k řešenému problému a co o něm vypovídají.

Snižování variability ovšem nesmí být obecným cílem. V biologii při šlechtitelství často dochází křížením různých druhů ke zvýšení kvality produktů a naopak rozmnožování probíhající v rámci úzkého výběru jedinců vede ke zvýšenému výskytu degenerativních jevů.

Modely statistického myšlení a jeho výuky lze zhruba rozdělit do dvou skupin. První jsou založeny na současných teoriích učení a rozvíjení duševní aktivity a pokoušejí se podat přímé návody k výuce statistiky nejen ve škole (Jones aj., 2000), ale i např. v řízení organizací (Hoerl a Snee, 2001). Druhá skupina modelů je založena na přímém pozorování výukového procesu či navrženého výukového či výzkumného projektu. Je tedy empirická a pokouší se objevit konstituční prvky statistického myšlení i rozpoznat úskalí, na které tento duševní proces naráží. Jisté podněty pro výukový proces jsou opět získány. Dva příklady modelů tohoto druhého typu uvedeme podrobněji.

První z nich je čtyřstupňový hierarchický model odvozený z výukové praxe. Ben-Zvi a Friedlander (1997) se v něm pokoušejí vysledovat různé etapy výchovy ke statistickému myšlení sledováním interakce mezi studenty (věk 13 až 15 let) a výpočetní technikou při zpracování předem zadaných dat,

tj. postupné zlepšování schopnosti řešit zadaný problém s podporou počítače. Takto byly identifikovány čtyři stupně:

Stupeň 0 (nekritické myšlení): soustředění na extrémní rysy dat, práce v základním programovém režimu (default options: používání lineárních stupnic i při velkém rozsahu nerovnoměrně rozložených dat, nevhodná volba intervalů u histogramů atd.), volba grafických reprezentací podle estetických měřítek a jejich neúčelná nadprodukce znemožňující dospět k jednoduchým shrnujícím závěrům.

1. stupeň (smysluplné zdůvodnění zvolené reprezentace): grafické či numerické reprezentace respektují data a způsob jejich získání, schopnost demonstrovat výhodnost zvoleného přístupu předvedením jiných méně vhodných reprezentací. Tzv. PCAI cyklus (Pose, Collect, Analyze, Interpret), tj. formuluj problém, shromáždí či uprav data, analyzuj a interpretuj, je ještě málo rozvinutý a převažuje grafická reprezentace.

2. stupeň (smysluplné používání různých reprezentací): schopnost prezentovat data různým způsobem (graficky i numericky) s použitím standardních postupů.

3. stupeň (tvůrčí myšlení): schopnost vymýšlet nestandardní reprezentace dat vhodné pro řešení problémů.

Konkrétní řešené příklady a rozbor jednotlivých stadií jsou v citované práci Ben-Zvi a Friedlandera. Model svou hierarchickou stavbou slouží učitelům jako popis cesty, kterou jeho žáci při kultivaci statistického myšlení mají projít, a jako návod, jak jim pomáhat. Jeho předností je, že zahrnuje počítačové zpracování jako neodmyslitelnou součást moderní statistiky a zdůrazňuje široké a snadno dostupné možnosti grafické reprezentace. Získání zkušeností při jejím vytváření je pro statistickou gramotnost (viz Závěr) důležité proto, že mnoho statistických výsledků je v odborné literatuře i v médiích graficky prezentováno nevhodnou a hrubě zjednodušující, často i záměrně zavádějící formou („sugestivní“ volbou měřítka či kombinací několika měřítek v jednom grafu, soustředěním na střední hodnoty bez představy o jejich variabilitě atd.). Zaměřením na zpracování zadaných dat však model Ben-Zvi a Friedlandera pomáhá pro každý problém klíčovou etapu volby vhodných proměnných a jejich sběr.

Druhý model odvozený z empirického zkoumání prezentuje konferenční příspěvek Wild a Pfannkuch (1999); vychází z dlouhodobé práce se čtyřmi skupinami účastníků (počet ve skupině 5 až 6) vybraných ze studentů kolem 20 let, profesionálů z oboru psychologie a statistiků (Pfannkuch, 1999). Výsledkem je představa o čtyřrozměrném myšlenkovém procesu, jehož všechny komponenty probíhají současně. První dimenze je výzkumný cyklus: získání dat, hodnocení a interpretace. Druhou dimenzi tvoří soubor úvah o kvalitě

dat, jejich optimální numerické reprezentaci a o zvýraznění variability. Třetí dimenze je vyšetřovací cyklus, jehož cílem je sladění představ o procesu se získanou informací včetně posouzení konsistence dat a možných nežádoucích vazeb (např. emotivní kontextuální přístup řešitele). Poslední etapa nazvaná autory *disposition*, tj. sestavení či povaha, je průběžnou úvahou nad problémem a jeho řešením s řadou složek, jako zájem o problém a jeho řešení, posouzení výsledku zdravým rozumem (pokud je to možné), představivost, skryté významy dat (data mohou vypovídat i o něčem, nač jsme se neptali) i závěrů.

Převážná část disertace Pfannkuch (1999) je přístupná na Internetu a obsahuje řadu podrobností, řešené příklady a záznamy diskusí nad nimi atd. Ty jsou patrně nejprínosnější a jsou otevřeny k dalším, autorkou neuvažovaným interpretacím. Model je analytickou studií myšlení v pravděpodobnostních a statistických souvislostech, což je jeho hlavní přínos. Volbou účastníků projektu je zaměřen na osoby již prošlé základním příp. i vysokoškolským vzděláním. Proto rozpoznává jeho nedostatky a naznačuje, co je v něm třeba zlepšit či změnit. Poněkud však pomíjí skutečnost, že současný student i profesionál – s výjimkou etapy shromažďování dat – *statisticky myslí* v úzké kooperaci s počítačem. Oba modely (Ben-Zvi, Friedlander a Wild, Pfannkuch) se tak vhodně doplňují.

Zvláštní kapitolou je pravděpodobnostní a statistické myšlení v geometrických souvislostech. Geometrickou pravděpodobnost a statistiku školní osnovy zcela pomíjejí a je obsažena jen okrajově i ve vysokoškolských kursech. Příčinou je patrně její pozdější vznik a dosti samostatný vývoj, probíhající dlouhou dobu především v rámci praktických aplikací (kvantifikace a interpretace obrazů materiálových a biologických struktur, geologický a zvláště důlní průzkum). Teorie náhodných systémů (tzv. procesů) geometrických objektů v eukleidovském prostoru, o nichž získáváme informace prostředky geometrického výběru (řezy, projekce) se objevuje až v sedmdesátých letech minulého století, problematika náhodných tvarů ještě později. Jejím většímu rozšíření brání nedostatečné znalosti z geometrie, jejíž výuka byla v minulém století postupně omezována (všude na světě). Přesto se s intuitivními (tj. opírajícími se o podvědomě zpracovanou dlouhodobou zkušenost) aplikacemi statistického myšlení s geometrickým obsahem setkáváme doslova na každém kroku a při každém pohledu kolem sebe. Umělé i přírodní objekty (předměty denní potřeby, potraviny, chovná zvířata) mají rozměry i tvary náhodné v poměrně úzkých rozmezích, v jejichž rámci je hodnotíme a přiřazujeme jim určité vlastnosti. V odborné praxi jsou nejběžnějšími příklady odhady vlastností materiálů z jejich krystalických struktur (tzv. velikost zrna a její variabilita u polykrystalů), rozpoznávání geometricky charakterizova-

telných patologických změn orgánů a tkání, hodnocení územního a správního členění s ohledem na účelnost a operativnost. Problematikou výuky a přechodu od intuitivního ke kvalifikovanému přístupu bylo zatím věnováno jen minimální množství prací. Výjimku představuje poměrně rozsáhlý průzkum dětského vnímání náhodnosti rovinného rozložení bodů (Green, 1992).

3. Poznatky z práce se žáky

Děti se s náhodou a rizikem setkávají již v předškolním věku při hrách, doma i v dětském kolektivu. Proto se u nich vyvíjí intuitivní chápání nejistoty některých dějů, v méně přesné podobě i jejich jistoty či naopak nemožnosti („maminka bude určitě doma“, „babička nemůže hrát kopanou“). Rozvoj tohoto intuitivního myšlení je třeba včas správným způsobem ovlivňovat vhodnými hrami a jejich výkladem. Way (1990) uvádí několik příkladů takových kolektivních her pro děti ve věku dvou až tří let. Děti se v nich učí diferencovat běžné předměty (na př. části svých oděvů a obutí) podle jejich vlastností (barva, hladkost atd.), seznamují se s problematikou stejně i nestejně možných jevů prostřednictvím her řízených hodem jedné či více mincí nebo tažením losů z urny.

Rozsáhlý průzkum přístupu žáků ve věku 6 až 14 let k náhodným jevům proběhl v Itálii v letech 1975 až 1980 (Perelli d'Argenzio aj., 1991). Jako první stupeň výchovy jsou doporučeny jazykové úlohy, neboť bylo zjištěno, že žáci mají potíže s termíny běžně užívanými ve statistice a pravděpodobnosti.

Do neúplných vět následujícího typu měli žáci doplňovat slova *jistý*, *nejistý*, *nemožný*:

- Je ..., že dnes je pondělí.
- Je ..., že na konci roku postoupím do vyšší třídy.
- Je ..., že naši fotbalisté vyhraji turnaj.
- Je ..., aby u nás v létě mrzlo.

Žáci na jedné straně nevěděli přesně, co znamená slovo nejisté, na druhé straně byly jejich odpovědi ovlivněny emocemi, takže úspěch svých fotbalistů vesměs považovali za jistý. Podobnými větami byly poté testovány různé stupně nejistoty: *velmi nejisté*, *dosti nejisté*, *trochu nejisté*.

Schopnost kvalitativně posoudit elementární pravděpodobnosti byla sledována na urnovém modelu s průhlednými nádobami a černými (Č) a bílými (B) koulemi s následujícím obsazením: (6B,2Č), (3Č), (5B,3Č), (4B). Žáci měli určit urnu, u níž je nejvyšší pravděpodobnost vytáhnutí bílé. Překvapivě se objevovaly odpovědi „1. urna“ (nejvíce bílých) a „je to jedno“! Po podrobném vysvětlení a diskusi byly řešeny složitější úlohy s obsazením (12B,4Č), (13B,10Č) vedoucí již k relativnímu porovnání četností úspěchu.

Jako kategoriální proměnné byly hodnoceny např. obliba předmětu (ano, ne) a frekvence měsíců narození žáků s otázkami typu: udá-li každý student dva oblíbené předměty (z deseti), kolik dostaneme údajů dohromady?, lze předměty srovnat podle obliby?, narodil se v každém měsíci aspoň jeden žák?, má každý žák svůj měsíc narození?

Na vhodných příkladech byla postupně zavedena střední hodnota (žáci mají různý počet čistých výkresů; jak je rozdělit, aby měl každý stejně) i medián.

Nejdůležitější poznatky z těchto i jiných studií jsou následující: řešení úlohy musejí být konkrétní a mají se žáků bezprostředně dotýkat, data si mají žáci připravovat sami, optimální je, když sami navrhnou, jaký typ dat potřebují pro řešení dané úlohy.

Příkladem statistických aktivit určených pro žáky středních škol ve Velké Británii jsou každoročně pořádané statistické soutěže o ceny (několik kategorií zahrnujících vždy dva sousední ročníky). Jedná se o žáky připravené dotazníkové akce na aktuální problémy (doprava, spotřeba energie atd.), dále o odborné práce (např. výživná hodnota školních svačin, lokální ekologie) a komerční výzkum (nabídka zboží, kvalita produktů). Pro správnou formulaci otázek v dotazníkové akci se používá termínu *umění dotazníku*; spočívá v citlivé volbě nerozporných a neduplicitních otázek, jejichž přítomnost komplikuje a zdražuje celou akci.

Poznatky zjištěné na nejnižším stupni jsou platné i pro stupně vyšší, pouze úlohy jsou obtížnější. Přímá účast studentů na přípravě dat je velmi důležitá. Populárním příkladem je úloha z knihy Scheaffer aj., (1996) - měření průměrů tenisových míčků; při ní obtížnost získání dat vede studenty ke kritickému přístupu k datům ostatním. Zároveň je učí chápat statistický pokus jako celek, spočívající ve formulaci problému, získání dat, jejich analýze a interpretaci - tedy již zmíněný „PCAI“ přístup, a nikoliv jako pouhé numerické cvičení s hotovými daty. Plánování celého pokusu je vždy důležité, zvláštní pozornost je mu však třeba věnovat ve všech případech, kdy existuje možnost ovlivnění pokusu emocemi jeho organizátorů, tj. snahou něco předem formulovaného výběrovým šetřením prokázat. Např. při výzkumu veřejného mínění ať již formou dotazu či dotazníkovou akcí jsou neutralita otázky a postup při tvorbě výběrového souboru zcela klíčové; sugestivní otázka spolu s nevhodnou volbou místa a času (diskriminující jisté skupiny obyvatelstva) výsledek zcela znehodnocují. Přínosem podrobné diskuse se žáky je nejen zajištění korektnosti pokusu, ale i skepse ve vztahu k akcím různých agentur. Žáci se naučí respektovat pravidlo, že neznáme-li přesnou formulaci otázky, nepřijímáme nekriticky výsledky výzkumu. Skeptický vztah k datům je důležitou složkou

statistického myšlení vedoucí především k rozpoznání nevhodné volby sledované proměnné (měřené veličiny, otázky) či jejího chybného zjišťování (Wild a Pfannkuch, 1999).

Diskuse učitele se žáky i žáků mezi sebou je obecně důležitější, než formálně správné, avšak do hloubky nepromyšlené řešení. Učebnicové příklady často nemají přesně formulované předpoklady, resp. předpokládají automatické využití právě probírané látky. Problematický pak může být např. předpoklad stejných pravděpodobností posuzovaných reálných jevů. V disertaci Pfannkuchové je diskutován následující příklad:

V rodině mají tři dcery. Jaké pohlaví očekáváte u dalšího dítěte?

Účastníci projektu s jedinou výjimkou odpovídali na základě analogie (v příkladu však explicitně nevyslovené) s vrhem mince, že obě pohlaví jsou stejně pravděpodobná. Výjimku představovala zdravotní sestra, která byla na základě své zkušenosti přesvědčená, že v podobných případech je pravděpodobnost dalšího děvčete podstatně vyšší než $\frac{1}{2}$. Svou subjektivní zkušenost uplatňovala i při řešení jiných příkladů a odlišovala se tak od ostatních členů skupiny. I když pomineme skutečnost, že v průměru je pravděpodobnost narození chlapce o něco vyšší než $\frac{1}{2}$, mezi 0.51 a 0.52, zkušenost s početnými rodinami činí příklad sporným. Chceme dostat formálně matematicky správné řešení nebo odpověď na problém z reálného života? Nemusíme v druhém případě respektovat pravidlo, že každá reálná „činnost“ je výsledkem posloupnosti celé řady procesů? Zároveň se však nesmí zapomínat na to, že subjektivní zkušenost může být zavádějící, protože je obvykle založena na nedostatečně velkém počtu jevů.

Dalším zajímavým poznatkem je, že zatímco u učebnicových úloh podobných výše diskutované mají žáci tendenci aplikovat představy o stejných pravděpodobnostech jevů, a z nich plynoucí variabilitu výsledků uvažovat, pro variabilitu reálných jevů se vesměs pokoušejí hledat kauzální vysvětlení. Podobný sklon je pozorovatelný zcela obecně. Na příklad i při běžné variabilitě lokálního výskytu určitého onemocnění je pozitivní odchylka zhusta považována za předzvěst epidemie či důkaz zhoršení životního prostředí; zvláště citlivá jsou media, pro něž takové interpretace jsou atraktivní.

Za zmínku stojí jeden z mála projektů věnovaných geometrické pravděpodobnosti. Proběhl na britských školách s počtem kolem 6000 žáků (Green, 1992) a jeho obsahem bylo zjistit představu žáků o rovnoměrně náhodném rozložení bodů v rovinné oblasti. Řešená úloha byla následující: čtverec byl rozdělen na 16 malých čtverečků a v nich bylo třemi různými způsoby umístěno 16 bodů (v jejich středech, po jednom náhodně v každém z nich a konečně nepravidelné shluky dvou až tří v některých z nich); žáci měli za úkol vybrat případ, který považují za náhodný s legendou, že se jedná o kapky

deště (žáci ve věku 7 až 11 let) či sněhové vločky (11 až 16 let) na čtvercových střešních taškách. Zhruba 60% žáků považovalo za náhodné rozmístění jedno z těch, kdy v každém čtverci je právě jeden bod, shluky připouštělo jen 15 až 25% žáků (ostatní nevěděli nebo zamítli všechny případy jako neodpovídající jejich představě). Výsledek závisel na kontextu (zda se jednalo o déšť či sníh) a na věku (horší výsledek u starších žáků) a v podstatě žáci žádného věku nebyli na problém připraveni, pojem „náhodný“ jim nebyl jasný a jeho chápání se s věkem nelepšilo. Z vlastních zkušeností mohu konstatovat, že i dospělí absolventi vysokých škol pozorující rovnoměrně náhodné rozmístění bývají překvapeni jak přítomností a početností shluků, tak velikostí prázdných oblastí. Variantou úlohy bylo rozmisťování bodů samotnými žáky; výsledek byl velmi podobný. Závěr je, že jen mírně narušenou pravidelnost intuitivně považujeme za rovnoměrnou náhodnost či ekvivalentně za vzájemnou nezávislost poloh.

4. Závěr

V řadě citovaných prací včetně názvu amerického programu rozvoje statistické výuky je používán termín *literacy*, tj. *gramotnost*. Rozumí se pod ním seznámení s „metodami sběru, úpravy a interpretace dat, jejich grafickým znázorňováním, testováním a prezentací závěrů z nich vyvozených“. Patrně lze přijmout tvrzení, že se jedná o specifickou gramotnost informačního věku, a připustit skutečnost, jejíž následky vidíme velmi často kolem sebe, že naše školy opouštějí absolventi v tomto slova smyslu *negramotní*. Potom je však nezbytným důsledkem, že této gramotnosti bychom měli žáky učit po celou dobu jejich vzdělávání a se stejnou péčí, s jakou je učíme číst, psát a vyjadřovat se. Učit je významu slov v souvislosti s náhodnými ději používanými, vnímání omezené rozmanitosti tvarů a rozměrů typově stejných objektů, seznámit je alespoň se základními typy rozdělení (hrubá představa pouze o symetrickém normálním a rovnoměrném rozdělení nestačí) a s pojmy závislosti a nezávislosti příp. korelace dat, nevynechat vícerozměrné náhodné veličiny. A konečně je důležité při výuce nezapomínat, že žáci už leccos z toho, co je učíme, intuitivně znají. Jinak hrozí nebezpečí, že životem již ověřené zkušenosti nahradíme formálními „školskými“ pojmy. S učením lze snadno začít od první třídy; první napsaná písmenka a jejich postupně nabývané a zručnosti úměrné přibližování žádanému vzoru či rychlost a správnost čtení budou ideálním příkladem náhodného procesu, navíc časově závislého a k soutěživosti vyzývajícího. A které předměty vyšších ročníků nepojednávají o důsledcích náhodných jevů a nepracují s jejich charakteristikami? Na samém počátku dochované literatury, v indické Ridgvédě, najdeme náрек nešťastného hráče

v kostky, v Mahábháratě se setkáme s výběrovým šetřením, losem je v bibli dělena země mezi izraelské kmeny a talmud je plný pravděpodobnostních příkladů. A dějepis? Hérodotos odhaduje časové rozpětí mezi dvěma egyptskými vládci z počtu generací, které se vystřídají za sto let, v Thúkídidovi Platajští odhadují ze střední tloušťky cihel délku žebříků potřebnou k překonání hradeb, do nichž je uzavřel nepřítel, a všechny šťastné a nešťastné bitvy naší historie jsou do značné míry důsledkem souhry mnoha současně probíhajících a vzájemně souvisejících náhodných procesů. Jak málo vzájemných srovnání států obsahují učebnice zeměpisu a kolik by se z něj studenti naučili při jejich sestavování a grafickém znázorňování. Jednoduchá stochastická analýza textů prováděná studenty při jazykové výuce by zcela určitě přispěla k jejich snaze o zlepšení, která by pak měla charakter závodu s čísly. A kolik podnětů statistice přinesla biologie a sociologie, jak lze bez ní vykládat pravidla genetického výběru? Aplikace ve fyzice snad ani není třeba zmiňovat. Jestliže tento společný základ tolika jevů ve vzdělávání pomíjíme, jak můžeme být spokojeni?

Můžeme se také zamyslet nad výukou matematice. V ní jedno plyne z druhého a skoro všechno souvisí se vším. Její význam pro rozvoj exaktního myšlení žáků je nesporný, zachycuje mnoho z historie poznání reálného světa, byla a je nezbytná pro rozvoj techniky, přírodních i společenských věd. Ale kolik žáků ve svém životě po ukončení školy vyřeší alespoň jednu kvadratickou rovnici, zkonstruuje si trojúhelník ze tří prvků a zkoumá rozklad čísla na prvočinitele ve snaze prokázat, že se jedná prvočíslo? Naproti tomu pravděpodobnostní a statistické myšlení od něj bude požadováno celý život. Jeho výuka by bez ohledu na osnovy měla být průběžnou snahou všech učitelů matematiky od první do poslední třídy. Měli by se snažit do ní zapojit i učitele ostatních předmětů, resp. jim v ní pomoci upozorněním, kde v jejich předmětech lze či je vhodné tento typ myšlení uplatnit. Internet je v současné době naprosto nevyčerpatelným zdrojem podnětů a vědomostí, příkladů a interaktivních cvičení. Zvláště podnětný je *on line* časopis *Journal of Statistics Education*; v Dodatku je pro ilustraci uveden obsah jeho posledního 11. ročníku včetně internetové adresy. V něm lze nalézt i řadu citovaných prací.

Závěrem uvedme shrnutí dosavadního stavu poznání o statistickém myšlení podané Chance (2002) vymezením jeho základních šesti etap:

1. Hledání vhodného typu dat a způsobu jejich získání pro daný problém.
2. Určení vhodných proměnných.
3. Chápání procesu od získání dat po jeho závěr jako jediného celku.
4. Neustálá skepse ve vztahu k postupu i výsledku.

5. Úvaha o vztahu dat k problému, interpretace výsledku v nestatistické terminologii.
6. Odchýlení od učebnicového postupu, je-li účelné.

Standardně používaný postup, kdy studenti mají za úkol na zadaných datech provést vyhodnocení základních charakteristik polohy a variability skutečné či modelové náhodné proměnné, příp. provést test předložené hypotézy, pouze cvičí výpočetní techniku, avšak narušuje prakticky všechna tato pravidla.

Literatura

- [1] Ben-Zvi D. a Friedlander A. (1997). *Statistical thinking in a technological environment*. In: J. Garfield a G. Burrill (Eds.), *Research on the role of technology in teaching and learning statistics*. Voorburg, The Netherlands: International Statistical Institute, 45–55.
- [2] Box G.E.P., Hunter W.G. a Hunter J.A. (1978). *Statistics for Experimenters*. John Wiley and Sons, New York.
- [3] Burrill G. (1991). *Quantitative literacy in the United States*. In: R. Morris (ed.): *The teaching of statistics. Studies in mathematics education, Vol. 7.*, UNESCO, Paris, 81–93.
- [4] Chance B.L. (2002). *Components of Statistical Thinking and Implications for Instruction and Assessment*. Journal of Statistics Education 10 (2002), No. 3.
- [5] *Glossary of Statistical Terms* (1996). American Society for Quality (1996), Milwaukee, WI.
- [6] Green D. (1991). *School pupils' understanding of randomness*. In: R. Morris (ed.): *The teaching of statistics. Studies in mathematics education, Vol. 7.* UNESCO, Paris, 27–39.
- [7] Hawkins A. (1991). *The annual United Kingdom Statistics Prize*. In: R. Morris (ed.): *The teaching of statistics. Studies in mathematics education, Vol. 7.* UNESCO, Paris, 217–227.
- [8] Hoerl R.W. a Snee R.D. (2001). *Statistical thinking: Improving business performance*. Pacific Grove, CA: Duxbury.
- [9] Jones G., Thornton C., Langrall C., Mooney E., Perry B. a Putt I. (2000). *A framework for characterizing children's statistical thinking*. Mathematical Thinking and Learning, 2(4), 269–307.
- [10] Mallows C. (1998). “*The Zeroth Problem*”. The American Statistician, 52, 1–9.

- [11] Moore D.S. (1990). “*Uncertainty*”. In: L. A. Steen (ed.): *On the Shoulders of Giants*, National Academy Press, 95-173.
- [12] Ottaviani M.G. (1991). *A history of the teaching statistics in higher education in Europe and in the United states 1660 to 1915*. In: R. Morris (ed.): *The teaching of statistics. Studies in mathematics education*, Vol. 7. UNESCO, Paris, 243–252.
- [13] Perelli d’Annunzio M.P., Cutillo E.A. a Pesarin F. (1991). *The teaching of probability and statistics in italian compulsory schools*. In: R. Morris (ed.): *The teaching of statistics. Studies in mathematics education*, Vol. 7. UNESCO, Paris, 228–241.
- [14] Pfannkuch M. (1999). *Characteristics of Statistical Thinking in Empirical Inquiry*. Doctoral Thesis. The University of Auckland.
- [15] Pfannkuch M. a Wild C. (2002). *Statistical thinking models*. In: Brian Phillips (ed.): *ICOTS 6, Proceedings of the 6th International Conference on Teaching Statistics*.
- [16] Scheaffer R., Gnanadesikan M., Watkins A. a Witmer J. (1996). *Activity-Based Statistics*. Springer-Verlag Publishers, New York.
- [17] Snee R.D. (1990). *Statistical Thinking and Its Contribution to Total Quality*. The American Statistician, 44, 116–121.
- [18] Snee R.D. (1999). *Discussion: Development and Use of Statistical Thinking: A New Era*. International Statistical Review, 67, 255–258.
- [19] Way J. (1997). *Laying foundations in chance and data*. Reflections 21, No. 1., 223–265.
- [20] Wild C.J. a Pfannkuch M. (1999). *Statistical Thinking in Empirical Enquiry*. International Statistical Review, 67, 223–265.
- [21] Závodský P. (1992). *Vývoj statistické teorie na území Československa do roku 1848*. Federální statistický úřad a INFOSTAT, Praha.

JOURNAL OF STATISTICS EDUCATION

An International Journal on the Teaching and Learning of Statistics

electronická verze: <http://www.amstat.org/publications/jse/>

Volume 11, Number 3 (November 2003), Table of Contents

- C.M. Anderson-Cook a Sundar Dorai-Raj, *Making the Concepts of Power and Sample Size Relevant and Accessible to Students in Introductory Statistics Courses using Applets*.
- Jessica Utts, Barbara Sommer, Curt Acredolo, Michael W. Maher a Harry M. Matthews, *A Study Comparing Traditional and Hybrid Internet-Based Instruction in Introductory Statistics Classes*.

- David M. Reineke, Jeffrey Baggett a Abdulaziz Elfessi, *A Note on the Effect of Skewness, Kurtosis, and Shifting on One-Sample t and Sign Tests.*
- Steve H.K. Ng, *Studying the Effect of the Parameters of an Attribute Acceptance Sampling Plan on its Operating Characteristic Curve: A Visual Approach.*
- Pierre Duchesne, *Estimation of a Proportion with Survey Data.*
- Nyaradzo Mvududu, *A Cross-Cultural Study of the Connection Between Students' Attitudes Toward Statistics and the Use of Constructivist Strategies in the Course.*

Volume 11, Number 2 (July 2003), Table of Contents

- Michael A. Martin, *"It's like...you know": The Use of Analogies and Heuristics in Teaching Introductory Statistical Methods.*
- Paul J. Roback, *Teaching an Advanced Methods Course to a Mixed Audience.*
- Jonathan R.D. Kuhn, *Graphing Calculator Programs for Instructional Data Diagnostics and Statistical Inference.*
- Larry Feldman a Fred Morgan, *The Pedagogy and Probability of the Dice Game HOG.*
- Peter P. Howley, *Teaching How to Calibrate a Process Using Experimental Design and Analysis: The Ballistat.*
- Grete Heinz, Louis J. Peterson, Roger W. Johnson a Carter J. Kerk, *Exploring Relationships in Body Dimensions.*

Volume 11, Number 1 (March 2003), Table of Contents

- Daniel W. Schafer a Fred L. Ramsey, *Teaching the Craft of Data Analysis.*
- Mary Richardson a Byron Gajewski, *Archaeological Sampling Strategies.*
- Arthur Bakker, *The Early History of Average Values and Implications for Education.*
- Ian McLeod, Ying Zhang a Hao Yu, *Multiple-Choice Randomization.*
- Timothy S. Vaughan, *Teaching Statistical Concepts With Student-Specific Datasets.*
- Rose Martinez-Dawson, *Incorporating Laboratory Experiments in an Introductory Statistics Course.*
- James V. Pinto, Pin Ng a David S. Allen, *Logical Extremes, Beta, and the Power of the Test.*

ZÁKLADY PRAVDĚPODOBNOSTI

Josef Štěpán

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

josef.stepan@mff.cuni.cz

1. Co je to pravděpodobnost

Začneme matematickým modelem pro popis náhodných jevů a jejich pravděpodobností. Uvědomme si, že aniž bychom konstruovali konzistentní matematický model pro náhodu, intuitivně klademe na pravděpodobnost jisté „rozumné“ požadavky. Například aby pravděpodobnost toho, že na kostce padne pětka či šestka, byla součtem pravděpodobností pětky a šestky. Náš model takové představy zahrnuje a umožňuje nám provádět rigorózní analýzu fenoménu náhody.

Definice 1. Pravděpodobnostní prostor je trojice $(\Omega, \mathcal{F}, P)$, kde Ω je neprázdná množina, $\mathcal{F}$ některá algebra podmnožin Ω a P pravděpodobnostní množinová funkce (PMF) definovaná na algebře $\mathcal{F}$.

Specifikujeme: Algebra $\mathcal{F}$ je systém podmnožin Ω s vlastnostmi

$$\emptyset, \Omega \in \mathcal{F}, \quad F^c = \Omega \setminus F \in \mathcal{F}, \quad F \cap G, F \cup G \in \mathcal{F} \quad \text{pro } F, G \in \mathcal{F}. \quad (1)$$

PMF P je funkce definovaná na algebře $\mathcal{F}$ s vlastnostmi

$$\begin{aligned} P(\emptyset) &= 0, \quad P(\Omega) = 1, \quad 0 \leq P(F) \leq 1, \\ P(F \cup G) &= P(F) + P(G) \quad \text{pro } F \cap G = \emptyset. \end{aligned} \quad (2)$$

V kontextu náhodného pokusu interpretujeme trojici $(\Omega, \mathcal{F}, P)$ takto:

- Ω je seznam všech možných výsledků ω náhodného pokusu. Množinu Ω v konkrétních situacích volíme tak, aby elementární jevy ω byly těmi nejjemnějšími výsledky náhodného pokusu, které je třeba rozlišovat.

Klíčová slova: Pravděpodobnostní prostor, urnový model, charakteristiky náhodné veličiny, limitní chování, podmíněná pravděpodobnost, nezávislost

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839.

- Jednotlivé prvky $\omega \in \Omega$ se nazývají elementární jevy, tedy základní výsledky náhodného pokusu (například na kostce padne výsledek 1, 2, 3, 4, 5 nebo 6 – máme proto šest elementárních jevů).
- Množiny $F \in \mathcal{F}$ (náhodné jevy) jsou vlastnosti výsledku pokusu, kterým umíme přiřadit pravděpodobnost $P(F)$.
- $F \in \mathcal{F}$ s $P(F) = 1$ se nazývá jev jistý, je-li $P(F) = 0$, říkáme, že F je jev nemožný.

Povšimněme si, že vlastnosti $F \in \mathcal{F}$, tj. takové vlastnosti, kterým umíme přiřadit pravděpodobnost $P(F)$, vykazují stabilitu na standardní množinové operace. Dále si povšimněme, že PMF P měří pravděpodobnost náhodných jevů $F \in \mathcal{F}$ aditivně, tak jako měříme plochu nebo objem.

Vlastnosti (2) mají elementární důsledky, blíže viz věty 1.3, 1.4 v [1]:

$$P(F^c) = 1 - P(F), \quad (3)$$

$$P(F) \leq P(G), \quad P(G \setminus F) = P(G) - P(F) \quad \text{pro } F \subseteq G, \quad (4)$$

$$P\left(\bigcup_{k=1}^n F_k\right) = \sum_{1 \leq k \leq n} P(F_k) - \sum_{1 \leq k < j \leq n} P(F_k \cap F_j) + \dots + (-1)^{n+1} P(\cap_{k=1}^n F_k). \quad (5)$$

Speciálně

$$P(F \cup G) = P(F) + P(G) - P(F \cap G). \quad (6)$$

Velmi často je výsledek pokusu $\omega \in \Omega$ značně komplexní entita, zatímco nás zajímají jen některé jeho numerické vlastnosti, běžně označované $X_1(\omega), X_2(\omega), \dots$. Obrazem je následující definice.

Definice 2. $(\Omega, \mathcal{F}, P)$ buď pravděpodobnostní prostor. Reálná funkce X definovaná na Ω s vlastnostmi:

$$\text{množina jejích hodnot } X(\Omega) \text{ je konečná,} \quad (7)$$

$$\text{množina } [X = x] = \{\omega \in \Omega : X(\omega) = x\} \in \mathcal{F} \quad \text{pro každé } x \in \mathbb{R}, \quad (8)$$

se nazývá náhodná veličina (NV).

Všimněme si zejména, že požadujeme $[X = x] \in \mathcal{F}$, tedy $[X = x]$ je náhodný jev a jsme schopni říci, jakou má pravděpodobnost. Zřejmě $P[X = x]$ je pravděpodobnost toho, že NV X nabývá hodnoty x . Podobně $P[X \leq x]$

je pravděpodobnost toho, že NV X má hodnotu, která je menší nebo rovna číslu x .

Počet pravděpodobnosti znal ve svých počátcích ([1], dodatek A4) pouze kombinatorický pravděpodobnostní prostor, který budeme nyní definovat.

Definice 3. Ω buď konečná neprázdná množina, $\mathcal{F}$ algebra všech jejích podmnožin a PMF P buď definovaná jako

$$P(F) = \frac{|F|}{|\Omega|}, \quad \text{kde } F \in \mathcal{F} \text{ a } |\cdot| \text{ je počet prvků množiny.}$$

Trojice $(\Omega, \mathcal{F}, P)$ se nazývá kombinatorický pravděpodobnostní prostor nad Ω (KPP).

Které pokusy jsou modelovány pomocí KPP nad Ω ? Definice 3 říká, že $P(\omega) = P(\{\omega\}) = |\Omega|^{-1}$ pro každý výsledek pokusu $\omega \in \Omega$. Pravděpodobnosti všech výsledků jsou stejné, KPP nad Ω je vhodný model pro pokusy, které neposkytují důvod preferovat některé $\omega_1 \in \Omega$ před jiným $\omega_2 \in \Omega$. Poznamenejme, že je-li $(\Omega, \mathcal{F}, P)$ KPP nad Ω , pak každá reálná funkce X definovaná na Ω je náhodná veličina.

Uvedeme některé příklady náhodných pokusů, pro které je kombinatorický pravděpodobnostní prostor vhodným modelem. Čtenář možná ocení příležitost zopakovat si základy klasické kombinatoriky v dodatku A1 z [1].

Příklad 1 (Výběr s vrácením). V osudí je a koulí černých a b koulí bílých. Z osudí postupně vytáhneme dvě koule, prvou taženou kouli vrátíme do osudí před druhým tahem. Uvažte náhodné jevy

$$B_1 = [\text{v prvním tahu bílá koule}], \quad B_2 = [\text{v druhém tahu bílá koule}]$$

a vypočtěte $P(B_1)$ a $P(B_2)$.

Jak dospět k evidentním pravděpodobnostem

$$P(B_1) = P(B_2) = \frac{b}{a+b} \quad ? \quad (9)$$

Jako Ω se lživě nabízí množina všech dvoučlenných posloupností 0 a 1, kde například (1,1) označuje výsledek, kdy byla dvakrát tažena bílá koule, (0,1) výsledek, kdy v prvním tahu byla tažena koule černá a v druhém tahu bílá. Pokud je však $a < b$, máme oprávněný pocit, že pokus preferuje posloupnosti více zaplněné jednotkami a tudíž, že model KPP nad tímto Ω není adekvátní (bylo by $P(B_1) = P(B_2) = \frac{1}{2}$).

Tuto obtíž odstraníme tak, že každé z $a + b$ koulí uměle přidělíme vlastní identitu tím, že černé koule očísujeme čísla $1, 2, \dots, a$ a bílé čísla $a + 1, a + 2, \dots, a + b$. KPP nad množinou Ω všech dvoučlenných posloupností čísel $1, 2, \dots, a + b$ je již jistě správný model pro náš pokus, protože umělým očíslováním jsme výsledky pokusu zrovnoprávnili. Jest tedy

$$|\Omega| = (a + b)^2, \quad |B_1| = b(a + b), \quad |B_2| = (a + b)b$$

a výsledek (9) je ověřen.

Příklad 2 (Výběr bez vracení). Pokus je stejný jako v příkladu 1 s tím rozdílem, že prvá tažená koule se do osudí nevrací. Jaké jsou pravděpodobnosti $P(B_1)$ a $P(B_2)$ nyní?

Opakujeme konstrukci z příkladu 1 s tím, že (zřejmě) Ω je nyní množinou všech dvoučlenných posloupností různých prvků množiny $1, 2, \dots, a + b$. Tedy

$$|\Omega| = (a + b)(a + b - 1), \quad |B_1| = b(a + b - 1) \quad |B_2| = ab + b(b - 1)$$

a (kupodivu?) též dostáváme pravděpodobnosti $P(B_1)$ a $P(B_2)$ jako v (9). Oba příklady jsou velmi speciální v kontextu známých Pólyových urnových schémat (viz [1], 3.6).

Příklad 3. Maxwellův-Boltzmannův model ve statistické fyzice uvažuje n rozlišitelných částic a r disjunktních částí fázového prostoru (příhrádek). Všechna rozmístění částic do příhrádek jsou stejně možná. Uvažte náhodné veličiny $K_1, K_2, \dots, K_r$, které označují počty částic v příhrádkách $1, 2, \dots, r$. Určete $P[K_i = k]$ pro $0 \leq k \leq n$.

Naším jediným problémem je matematizovat pojem rozmístění. Úplný popis fyzikální situace spočívá v tom, že pořídíme adresář částic, tj. $(a_1, a_2, \dots, a_n)$, kde $1 \leq a_k \leq r$ je adresa (číslo příhrádky) částice $1 \leq k \leq n$. Správný model je tedy KPP nad množinou Ω všech posloupností čísel $1, 2, \dots, r$ délky n . Jest

$$|\Omega| = r^n, \quad |[K_i = k]| = |[K_1 = k]| = \binom{n}{k} (r - 1)^{n-k}$$

Vzhledem k symetrické povaze pokusu je jedno, kterou z příhrádek uvažujeme, bez újmy na obecnosti tedy zvolíme pro výpočet tu první. Odtud plyne první rovnost. Pak vybereme k částic do první příhrádky, zbývajících $n - k$ částic rozmístíme do zbývajících $r - 1$ příhrádek libovolně, získáme tak druhou rovnost. Jest tedy pro každé $k \in \{0, 1, \dots, n\}$

$$P[K_i = k] = \binom{n}{k} \frac{(r - 1)^{n-k}}{r^n} = \binom{n}{k} \left(\frac{1}{r}\right)^k \left(1 - \frac{1}{r}\right)^{n-k}. \quad (10)$$

Více o Maxwellově-Boltzmannově modelu a jiných modelech statistické fyziky naleznete v [1], kapitoly 3.3, 3.4 a 3.5. Některé úlohy mohou být obtížné i z hlediska kombinatorického.

Příklad 4. Šestkrátě hodíme symetrickou kostkou. Uvážíme náhodné veličiny $K_1, K_2, \dots, K_6$, kde K_i označuje počet bodů dosažených v i -tém hodu. Vypočtěte pravděpodobnost p toho, že posloupnost $K_1, K_2, \dots, K_6$ je monotónní.

Modelem pro tento pokus je zřejmě množina $\Omega = \{1, 2, \dots, 6\}^6$ všech posloupností čísel $1, 2, \dots, 6$ délky 6 a KPP nad ní. Označíme

$$\begin{aligned} M_1 &= [K_1 \leq K_2 \leq \dots \leq K_6], & M_2 &= [K_1 \geq K_2 \geq \dots \geq K_6], \\ M_3 &= [K_1 = K_2 = \dots = K_6]. \end{aligned}$$

Protože počet neklesajících posloupností čísel $1, 2, \dots, n$ délky r je tolik co kombinací r -té třídy z n prvků s opakováním (viz dodatek A1, [1]), dostáváme

$$|\Omega| = 6^6, \quad |M_1| = \binom{6-1+6}{6} = \binom{11}{6}, \quad |M_3| = 6.$$

Použitím formule (6) obdržíme výsledek:

$$\begin{aligned} p &= P(M_1 \cup M_2) = P(M_1) + P(M_2) - P(M_1 \cap M_2) \\ &= 2P(M_1) - P(M_3) = 2 \frac{\binom{11}{6}}{6^6} - \frac{6}{6^6} \\ &\doteq 0,0197. \end{aligned}$$

Použili jsme formuli (6), její obecnější verze (5) umožní řešit úlohy jako je následující (viz [1], příklad 1.3).

Příklad 5. Počítač náhodně a rovnoměrně generuje permutace celých čísel $1, 2, \dots, n$ řádu n , $(k_1, k_2, \dots, k_n)$. Určete pravděpodobnost p_n toho, že bude generována permutace s alespoň jednou shodou, tj. taková permutace, že existuje $1 \leq j \leq n$ takové, že $k_j = j$.

K dosažení výsledku $\sum_{i=1}^n (-1)^{i+1} (n!)^{-1}$ použijte vzorec (5). Povšimněte si, že $\lim_{n \rightarrow \infty} p_n = 1 - e^{-1} \doteq 0,6321$.

Příklad 6. Uvažte znovu pokus z příkladu 4 a označme M maximum dosažených bodů, $M = \max\{K_1, K_2, \dots, K_6\}$. Určete pravděpodobnost $P[M = k]$ pro $1 \leq k \leq 6$.

Snadněji určíme pravděpodobnosti $P[M \leq k]$, neb je zřejmé, že jev $[M \leq k]$ zahrnuje právě k^6 jevů elementárních. Odsud, podle (4), je pro $1 \leq k \leq 6$

$$\begin{aligned} P[M = k] &= P([M \leq k] \setminus [M \leq k-1]) \\ &= P[M \leq k] - P[M \leq k-1] = \frac{k^6}{6^6} - \frac{(k-1)^6}{6^6}. \end{aligned} \quad (11)$$

Vypočteme-li pravděpodobnosti $p_k = P[M = k]$, dostaneme tabulku

k	1	2	3	4	5	6
p_k	0,00002	0,00135	0,01425	0,07217	0,24711	0,66510

Vidíme, že nejpravděpodobnější hodnota (modus) maxima M je 6 a jeho pravděpodobnostmi vážený průměr (střední hodnota) je $\sum_k k p_k = 5,5609$.

Naše příklady se dosud týkaly pouze experimentů s konečnou množinou výsledků Ω , které jsou považovány za stejně pravděpodobné a které jsou modelovány pomocí KPP nad Ω . Následující příklad vyžaduje model s nekonečnou množinou výsledků.

Příklad 7. Počítač generuje „náhodně a rovnoměrně“ čísla $\omega \in [0, 1]$. Určete pravděpodobnost p_n toho, že bude generováno číslo ω jehož dvojkový rozvoj má na n -tém místě jednotku (uvažujeme rozvoje s konečným počtem jedniček pro čísla typu 2^{-n}).

Jak modelovat tento náhodný pokus v rámci definice 1? Zřejmě musí být $\Omega = [0, 1]$. Je-li $X_n(\omega)$ n -tý člen dvojkového rozvoje čísla ω , pak $p_n = P[X_n = 1]$ a $(\Omega, \mathcal{F}, P)$ musí tedy být takový prostor, že funkce $X_1, X_2, \dots$ jsou náhodné veličiny. Protože

$$\begin{aligned} [X_1 = 1] &= \left[\frac{1}{2}, 1 \right], \quad [X_2 = 1] = \left[\frac{1}{4}, \frac{1}{2} \right) \cup \left[\frac{3}{4}, 1 \right], \dots, \\ [X_n = 1] &= \bigcup_{\substack{1 \leq k < 2^n \\ k \text{ sudé}}} \left[\frac{k-1}{2^n}, \frac{k}{2^n} \right) \cup \left[\frac{2^n-1}{2^n}, 1 \right], \dots, \end{aligned} \quad (12)$$

je tento požadavek splněn, když algebra $\mathcal{F}$ zahrnuje všechna konečná sjednocení disjunktních intervalů. Naštěstí systém těchto sjednocení je sám již algebrou a problém dvojice $(\Omega, \mathcal{F})$ je vyřešen. Jak definovat pravděpodobnost $P(F)$? Generátor vybírá čísla $\omega \in [0, 1]$ „rovnoměrně“, je tedy třeba,

aby pravděpodobnost intervalu I byla rovna jeho délce $|I|$. Abychom vyhověli požadavku na aditivitu pravděpodobnosti P , definujme pravděpodobnost sjednocení disjunktních intervalů jako součet jejich délek, tedy

$$P(F) = \sum_{k=1}^n |I_k|, \text{ kde } F = \bigcup_{k=1}^n I_k, \quad I_k \cap I_j = \emptyset \text{ pro } k \neq j. \quad (13)$$

Není úplně snadné se přesvědčit, že definice (13) je korektní, tj. že hodnota $P(F)$ nezávisí na volbě rozkladu $I_1, I_2, \dots, I_n$, a že P je PMF ve smyslu definice 1.

Umluvíme se, že takto konstruovaný pravděpodobnostní prostor bude nazýván generátor náhodných čísel v intervalu $[0, 1]$. Jeho konstrukce je taková, že každá funkce X_n , kde $X_n(\omega)$ je n -tý člen dvojkového rozvoje ω , je náhodná veličina. V rámci tohoto modelu dostáváme pomocí (12) očekávaný výsledek

$$p_n = P[X_n = 1] = 2^{n-1} \frac{1}{2^n} = \frac{1}{2}. \quad (14)$$

Pokusíme se zobecnit pravděpodobnostní chování náhodných veličin, které prozatím vstupovaly do našich příkladů.

Pro následující text se domluvme, že znak $X \sim R$ budeme číst *náhodná veličina X má rozdělení R* , nebo *X se řídí rozdělením R* .

Definice 4. Řekneme, že NV má alternativní rozdělení s parametrem $p \in [0, 1]$, píšeme $X \sim Alt(p)$, když X má pouze hodnoty 0 a 1 tak, že

$$P[X = 1] = p, \quad \text{a} \quad P[X = 0] = q = 1 - p.$$

Vrátíme se k příkladům 1 a 2, označíme $X_1 = I_{B_1}$ a $X_2 = I_{B_2}$ ($I_B(\omega)$ nabývá hodnot nula a jedna; $I_B(\omega)$ je jedna právě tehdy, když $\omega \in B$). X_1 a X_2 jsou tedy indikátory bílé barvy v prvním a druhém tahu. Zřejmě je

$$X_i \sim Alt\left(\frac{b}{a+b}\right) \quad \text{pro } i = 1, 2.$$

Také dvojkové souřadnice v příkladu 7 jsou takové, že $X_n \sim Alt\left(\frac{1}{2}\right)$.

Definice 5. Řekneme, že NV X má binomické rozdělení s parametry $n \in \mathbb{N}$ a $p \in [0, 1]$, píšeme $X \sim Bi(n, p)$, když X má hodnoty v množině $\{0, 1, \dots, n\}$ a platí, že

$$P[X = k] = \binom{n}{k} p^k q^{n-k}, \quad 0 \leq k \leq n, \quad q = 1 - p.$$

Poznamenejme, že aditivita P si vynucuje, aby platilo

$$\sum_{k=0}^n P[X = k] = P\left(\bigcup_{k=0}^n [X = k]\right) = P(\Omega) = 1,$$

což je podle binomické věty správná rovnost.

V příkladu 3 vystupují NV $K_1, K_2, \dots, K_r$, které označují počet částic v přihrádkách $1, 2, \dots, r$. V (10) jsme odvodili, že

$$K_j \sim Bi\left(n, \frac{1}{r}\right) \quad \text{pro } 1 \leq j \leq r, \quad (15)$$

je-li n celkový počet částic. Uvedeme další příklady NV s binomickým rozdělením.

Příklad 8. Mincí, jejíž rub je označen nulou a líc jednotkou, hodíme n -krát, S_n buď počet dosažených jednotek. Ukažte, že pak platí $S_n \sim Bi\left(n, \frac{1}{2}\right)$.

Pokus je modelován jako KPP nad množinou Ω všech nula-jednotkových posloupností délky n . Zřejmě jest pro $0 \leq k \leq n$

$$|\Omega| = 2^n, \quad |[S_n = k]| = \binom{n}{k}, \quad P[S_n = k] = \binom{n}{k} \left(\frac{1}{2}\right)^k \left(1 - \frac{1}{2}\right)^{n-k},$$

takže S_n má rozdělení $Bi\left(n, \frac{1}{2}\right)$.

Příklad 9. V osudí je a koulí černých a b koulí bílých. Z osudí postupně vybíráme n koulí, taženou kouli vždy vracíme. K_n buď NV, která označuje počet tažených koulí bílé barvy. Dokažte, že $K_n \sim Bi\left(n, \frac{b}{a+b}\right)$.

Tento pokus (viz příklad 1) modelujeme pomocí KPP nad množinou Ω všech posloupností čísel $1, 2, \dots, a+b$ délky n . Dostáváme

$$|\Omega| = (a+b)^n \quad \text{a} \quad |[K_n = k]| = \binom{n}{k} b^k a^{n-k},$$

takže $K_n \sim Bi\left(n, \frac{b}{a+b}\right)$.

Tyto definice samozřejmě nevyčerpávají všechny možnosti pravděpodobnostního chování náhodných veličin, ukázkou je maximum M v příkladu 6. Následující charakteristiky toto chování částečně popisují.

Definice 6. Je-li X náhodná veličina, pak čísla

$$\mathbf{E}X = \sum_x xP[X = x] \quad \text{a} \quad \text{var}X = \sum_x (x - \mathbf{E}X)^2 P[X = x]$$

nazýváme střední hodnota a rozptyl náhodné veličiny X .

Poznamenejme, že součty $\sum_x$ v předchozích formulích jsou součty konečné, protože $P[X = x] \neq 0$ pouze v konečně mnoha případech. Poznamenejme také, že interpretujeme-li $P[X = x]$ jako hmotu umístěnou do bodu x , pak $\mathbf{E}X$ je těžiště a $\text{var}X$ moment setrvačnosti takto vzniklé soustavy hmotných bodů.

Jednoduché vlastnosti střední hodnoty jsou:

Věta 1. *Budte X a Y NV, a, b, c reálná čísla. Pak*

$$\mathbf{E}(aX + bY + c) = a\mathbf{E}X + b\mathbf{E}Y + c, \quad (16)$$

$$\text{když } P[X \geq a] \geq P[Y \geq a] \text{ pro všechna } a, \text{ pak } \mathbf{E}X \geq \mathbf{E}Y. \quad (17)$$

Důkaz je snadný, (16) například dostaneme aplikací následujícího vzorce.

Lemma 1. *Jsou-li X a Y náhodné veličiny, $f(x, y)$ reálná funkce definovaná na $\mathbb{R}^2$, pak $Z = f(X, Y)$ je také náhodná veličina a*

$$\mathbf{E}Z = \sum_{(x,y)} f(x, y)P[X = x, Y = y]. \quad (18)$$

Je tedy

$$\begin{aligned} \mathbf{E}(X + Y) &= \sum_{(x,y)} (x + y)P[X = x, Y = y] \\ &= \sum_x x \sum_y P[X = x, Y = y] + \sum_y y \sum_x P[X = x, Y = y] \\ &= \sum_x xP[X = x] + \sum_y yP[Y = y] = \mathbf{E}X + \mathbf{E}Y \end{aligned} \quad (19)$$

a také

$$\mathbf{E}X^2 = \sum_x x^2 P[X = x]. \quad (20)$$

Spočteme (18): jest

$$[Z = z] = \bigcup_{(x,y) \in A_z} [X = x, Y = y] \in \mathcal{F}, \quad A_z = \{(x, y) : f(x, y) = z\}.$$

Z je tedy náhodná veličina a aditivita P říká, že

$$P[Z = z] = \sum_{(x,y) \in A_z} P[X = x, Y = y].$$

Jest tudíž

$$\begin{aligned} \mathbf{E}Z &= \sum_z zP[Z = z] = \sum_z z \sum_{(x,y) \in A_z} P[X = x, Y = y] \\ &= \sum_{(x,y)} f(x, y)P[X = x, Y = y]. \end{aligned}$$

□

Věta 1 společně s právě dokázaným lemmatem verifikují i následující vlastnosti rozptylu.

Věta 2. Je-li X NV, a, b, c reálná čísla, pak

$$\text{var}X = \mathbf{E}(X - \mathbf{E}X)^2 = \mathbf{E}X^2 - (\mathbf{E}X)^2, \quad (21)$$

$$\text{var}(aX + b) = a^2 \text{var}X. \quad (22)$$

Poznámka 1. Formule (21) nás navádí, abychom počítali rozptyl $\text{var}X$ jako

$$\sum x^2 P[X = x] - \left(\sum x P[X = x] \right)^2.$$

Určete takto rozptyl maxima M v příkladu 6 (spočítali jsme $\mathbf{E}M = 5,56029$). Takto také snadno ověříte, že

$$\text{pro } X \sim \text{Alt}(p) \quad \text{je} \quad \mathbf{E}X = p \quad \text{a} \quad \text{var}X = pq \quad (23)$$

$$\text{pro } X \sim \text{Bi}(n, p) \quad \text{je} \quad \mathbf{E}X = np \quad \text{a} \quad \text{var}X = npq. \quad (24)$$

Uvažte ještě náhodnou veličinu X , která má rovnoměrné rozdělení na množině $\{x_1, x_2, \dots, x_n\}$, tj. $P[X = x_k] = 1/n$ pro $1 \leq k \leq n$. Střední hodnota je tedy v tomto případě aritmetický průměr

$$\mathbf{E}X = \bar{x} = \frac{1}{n} \sum_k x_k$$

a rozptyl je aritmetický průměr čtverců odchylek od $\bar{x}$

$$\text{var}X = s^2 = \frac{1}{n} \sum_k (x_k - \bar{x})^2.$$

Zvolte $x_k = k$, určete $\bar{x}$ a s^2 .

Skutečný význam rozptylu ukazuje následující nerovnost.

Věta 3 (Čebyševova nerovnost I). Je-li X náhodná veličina a $\varepsilon > 0$, pak

$$P[|X - \mathbf{E}X| \geq \varepsilon] \leq \frac{\text{var}X}{\varepsilon^2}.$$

Čebyševova nerovnost říká, že s klesajícím rozptylem se zvětšuje koncentrace pravděpodobnosti v libovolném okolí střední hodnoty.

Důkaz je snadný. Je-li Y funkce přiřazující hodnotu 1 náhodnému jevu $[|X - \mathbf{E}X| \geq \varepsilon]$ a hodnotu 0 opačnému jevu, pak $Y \sim \text{Alt}(p)$, kde $p = P[|X - \mathbf{E}X| \geq \varepsilon]$. Podle (23) je

$$P[|X - \mathbf{E}X| \geq \varepsilon] = \mathbf{E}Y \leq \mathbf{E}[\varepsilon^{-2}(X - \mathbf{E}X)^2] = \varepsilon^{-2}\text{var}X,$$

kde prvá nerovnost plyne z (17), protože (!) $Y \leq \varepsilon^{-2}(X - \mathbf{E}X)^2$ a druhá rovnost je důsledek linearit (16).

Pomocí $\mathbf{E}X = \mu$ a $\text{var}X = \sigma^2$ můžeme konstruovat interval, který pokrývá hodnoty X s velkou pravděpodobností $1 - \alpha$ (třeba $\alpha = 0,05$). V Čebyševově nerovnosti volte $\varepsilon = \frac{\sigma}{\sqrt{\alpha}}$ a dostanete

$$P\left[\mu - \frac{\sigma}{\sqrt{\alpha}} \leq X \leq \mu + \frac{\sigma}{\sqrt{\alpha}}\right] \geq 1 - \frac{\sigma^2}{\varepsilon^2} = 1 - \alpha.$$

Pro $\alpha = 0,05$ je $\frac{1}{\sqrt{\alpha}} \doteq 4,5$ a dostáváme

$$P[\mu - 4,5\sigma \leq X \leq \mu + 4,5\sigma] \geq 0,95. \quad (25)$$

Tento odhad, díky své univerzalitě, příliš užitečný není. Dodatečná informace o typu rozdělení NV X může interval $[\mu - 4,5\sigma, \mu + 4,5\sigma]$ při zachování nerovnosti (25) podstatně zkrátit, jak ukážeme v závěru tohoto odstavce.

Vyšetříme limitní chování binomických pravděpodobností $Bi(n, p)$ ve dvou zcela odlišných situacích.

A. $Bi(n, p)$ pro velké hodnoty parametru n a malé hodnoty pravděpodobnosti p , je-li $np = \lambda \ll n$.

Přesněji, předpokládáme-li, že $\lim_{n \rightarrow \infty} np_n = \lambda > 0$, pak

$$\binom{n}{k} p_n^k (1 - p_n)^{n-k} = \frac{1}{k!} \left(1 - \frac{np_n}{n}\right)^n \frac{p_n(n-1)p_n \cdots (n-k+1)p_n}{(1-p_n)^k}$$

má při $n \rightarrow \infty$ limitu rovnou číslu $\frac{e^{-\lambda} \lambda^k}{k!}$ pro každé pevné $k = 0, 1, 2, \dots$

Úmluva 1. Řekneme, že náhodná veličina X má Poissonovo rozdělení $Po(\lambda)$, když nabývá hodnot $k = 0, 1, 2, \dots$ s pravděpodobnostmi

$$P[X = k] = \frac{e^{-\lambda} \lambda^k}{k!}.$$

Povšimneme si, že celková pravděpodobnost je

$$\sum_{k=0}^{\infty} P[X = k] = e^{-\lambda} \sum_{k=0}^{\infty} \frac{\lambda^k}{k!} = 1.$$

Definujme střední hodnotu i pro náhodnou veličinu $X \sim Po(\lambda)$ s nekonečně mnoha hodnotami. Je přirozené zobecnit definici 6 na nekonečné součty

$$\mathbf{E}X = \sum_{k=0}^{\infty} kP[X = k], \quad \text{var}X = \sum_{k=0}^{\infty} (k - \mathbf{E}X)^2 P[X = k].$$

Přímým výpočtem zjistíme, že $\mathbf{E}X = \text{var}X = \lambda$.

Uvědomíme si, že na tomto místě skutečně jde o úmluvu, v našem kontextu jsou náhodné veličiny funkce, které nabývají pouze konečně mnoha hodnot. Obecné zavedení střední hodnoty (a rozptylu) vyžaduje jistou míru opatrnosti při zacházení s nekonečným součtem či integrálem.

Výše uvedený limitní výpočet má nyní tvar následujícího tvrzení.

Věta 4 (Poissonova věta). *Uvažujme náhodné veličiny $X_n \sim Bi(n, p_n)$ takové, že $\lim_{n \rightarrow \infty} np_n = \lambda > 0$ a náhodnou veličinu Y , která má Poissonovo rozdělení $Po(\lambda)$. Pak*

$$\lim_{n \rightarrow \infty} P[X_n = k] = P[Y = k] \quad \text{pro} \quad k = 0, 1, 2, \dots$$

Poučení 1. Je-li $X_n \sim Bi(n, p)$, $n \rightarrow \infty$ a $p \rightarrow 0$, použijeme aproximaci $X_n \sim Po(np)$.

Podle vzorce (15) mají počty částic $K_1, K_2, \dots, K_r$ z příkladu 3 binomické rozdělení $Bi(n, 1/r)$. Při $n = 500$ a $r = 365$ můžeme bezpečně nahradit binomické rozdělení $Bi(500, 1/365)$ Poissonovým rozdělením $Po(500/365)$. Následující tabulka uvádí přesnou hodnotu $p_k = P[K_1 = k]$ v řádku druhém a její Poissonovu aproximaci a_k v řádku třetím.

k	0	1	2	3	4	5	6
p_k	0,2537	0,3484	0,2388	0,1089	0,0372	0,0101	0,0023
a_k	0,2541	0,3481	0,2385	0,1089	0,0373	0,0102	0,0023

Pravděpodobnost p_0 lze interpretovat jako pravděpodobnost toho, že ve skupině pěti set osob nemá nikdo narozeniny 1. ledna. Za jakých okolností a proč?

B. Druhým limitním chováním binomického rozdělení je situace $Bi(n, p)$ pro velké hodnoty n a „nikoliv příliš malé či velké“ hodnoty p .

Do této asymptotiky „magicky“ vstupuje hustota $\varphi(\cdot)$ a distribuční funkce $\Phi(\cdot)$ normálního rozdělení pravděpodobností, kterému také říkáme Gaussovo, tj.

$$\varphi(t) = \frac{1}{\sqrt{2\pi}} e^{-\frac{1}{2}t^2} \quad \text{a} \quad \Phi(x) = \int_{-\infty}^x \varphi(t) dt.$$

Uvědomme si nyní o jaký typ náhodné veličiny jde. V předchozí tabulce si můžeme povšimnout, že hodnoty 0, 1, až 5 dávají dohromady pravděpodobnost přes 99%, tedy s téměř jistotou lze tvrdit, že výsledek (počet lidí ze skupiny 500 osob, které mají narozeniny 1. ledna – za předpokladu, že rok má 365 dní, lidé se rodí rovnoměrně a ve skupině se nevyskytují dvojčata) bude jedna ze šesti hodnot (se stále vysokou pravděpodobností se lze omezit jen na hodnoty 0, 1, 2, 3).

Nyní je situace úplně jiná. Uvažujme velký počet náhodných pokusů ve kterých vystupuje náhodná veličina s alternativním rozdělením s dostatečně velkým parametrem p . Například při 600 hodech kostkou je za předpokladu, že šestka padne s pravděpodobností $1/6$, nejpravděpodobnější výsledek 100 šestek. Tato hodnota má ale pravděpodobnost zcela zanedbatelnou, menší než 5 %! Jmenovitě jde o hodnotu

$$P[X = 100] = \binom{600}{100} \frac{5^{500}}{6^{600}} \approx \frac{\sqrt{6}}{10\sqrt{10\pi}},$$

jak se lze přesvědčit použitím Stirlingova vzorce pro přibližný výpočet faktoriálu

$$n! \approx \sqrt{2\pi n} n^n e^{-n}.$$

Pro pevnou hodnotu p a vzrůstající počet pokusů roste i počet výsledků, které lze očekávat. Výsledkem je, že limitní rozdělení musí obsahovat nekonečně mnoho hodnot. Jednotlivé hodnoty ale mají nulovou pravděpodobnost a smysl má hovořit pouze o pravděpodobnosti nějakého intervalu. Zavedme si proto normální rozdělení.

Úmluva 2. Řekneme, že náhodná veličina Y má normální (Gaussovo) rozdělení $N(0, 1)$, když

$$P[a \leq Y \leq b] = \frac{1}{\sqrt{2\pi}} \int_a^b e^{-\frac{1}{2}t^2} dt \quad \text{pro} \quad -\infty \leq a \leq b \leq \infty.$$

Povšimneme si, že celková pravděpodobnost je

$$P[-\infty < Y < \infty] = \int_{-\infty}^{\infty} \varphi(t) dt = 1 \quad \text{a} \quad P[Y = y] = 0.$$

Definujme $\mathbf{E}Y = \int_{-\infty}^{\infty} t \varphi(t) dt$ a $\text{var}Y = \int_{-\infty}^{\infty} (t - \mathbf{E}Y)^2 \varphi(t) dt$.

Lehce spočítáme, že $\mathbf{E}Y = 0$ a $\text{var}Y = 1$ (proto $N(0, 1)$).

Ani tato úmluva nemůže být považována za definici. V našem kontextu Y není náhodná veličina.

Jakou souvislost má Gaussovo rozdělení s rozdělením binomickým? Pokušíme se „rozložit“ $P[a \leq Y \leq b]$ do binomických pravděpodobností při $p = \frac{1}{2}$. Budeme potřebovat několik poznatků z komplexní analýzy a základního kalkulu. Výsledky následujících výpočtů jsou shrnuty ve větách 5 a 6; čtenář, kterého nezajímá jejich odvození, proto může následující část vynechat.

Položíme

$$x_{nk} = \frac{k - np}{\sqrt{npq}} = \frac{2k - n}{\sqrt{n}}$$

a nejprve nahradíme integrál $\int_a^b \varphi(t) dt$ Riemannovou sumou:

$$\sum_k \frac{2}{\sqrt{2\pi n}} e^{-\frac{1}{2}x_{nk}^2} = \frac{1}{\pi\sqrt{n}} \sum_k \int_{-\infty}^{\infty} e^{-iux_{nk}} e^{-\frac{1}{2}u^2} du,$$

kde sčítáme přes všechna $0 \leq k \leq n$ taková, že $a \leq x_{nk} \leq b$, a používáme formuli

$$\varphi(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \cos(ux) e^{-\frac{1}{2}u^2} du = \frac{1}{2\pi} \int_{-\infty}^{\infty} e^{-iux} e^{-\frac{1}{2}u^2} du.$$

Nahradíme-li $\int_{-\infty}^{\infty}$ integrálem $\int_{-\frac{\pi\sqrt{n}}{2}}^{\frac{\pi\sqrt{n}}{2}}$ a použijeme-li standardní aproximaci exponenciely, je

$$\begin{aligned} e^{-\frac{1}{2}u^2} &\doteq \left(1 - \frac{u^2}{2n} + o(n^{-1})\right)^n = \cos\left(\frac{u}{\sqrt{n}}\right)^n = 2^{-n} \left(e^{i\frac{u}{\sqrt{n}}} + e^{-i\frac{u}{\sqrt{n}}}\right)^n \\ &= 2^{-n} \left(e^{-i\frac{u}{\sqrt{n}}}\right)^n \left(e^{2i\frac{u}{\sqrt{n}}} + 1\right)^n. \end{aligned}$$

Celkem dostáváme

$$P[a \leq Y \leq b] \doteq 2^{-n} \frac{1}{\pi\sqrt{n}} \sum_k \int_{-\pi\frac{\sqrt{n}}{2}}^{\pi\frac{\sqrt{n}}{2}} e^{\frac{-2iuk}{\sqrt{n}}} \left(e^{\frac{2iu}{\sqrt{n}}} + 1\right)^n du$$

a konečně, po substituci $t = \frac{2u}{\sqrt{n}}$, s využitím toho, že $\frac{1}{2\pi} \int_{-\pi}^{\pi} e^{-it(k-j)} dt$ je Kroneckerovo δ_{kj} , vypočítáme

$$\begin{aligned} P[a \leq Y \leq b] &\doteq 2^{-n} \frac{1}{2\pi} \sum_k \int_{-\pi}^{\pi} e^{-itk} (e^{it} + 1)^n dt \\ &= \sum_{k: a \leq x_{nk} \leq b} \binom{n}{k} 2^{-n} = P \left[a \leq \frac{2X_n - n}{\sqrt{n}} \leq b \right] \\ &= P \left[a \leq \frac{X_n - np}{\sqrt{npq}} \leq b \right], \end{aligned}$$

kde $X_n \sim Bi(n, \frac{1}{2})$.

Tuto heuristiku lze realizovat korektně pro každé pevné $p \in (0, 1)$ (viz [1], odstavec 4.5). Platí

Věta 5 (Moivreova-Laplaceova věta lokální). Pro $p \in (0, 1)$ a stejnoměrně pro $0 \leq k \leq n$ platí:

$$\binom{n}{k} p^k q^{n-k} = \frac{1}{\sqrt{npq}} \varphi(x_{nk}) + o(n^{-\frac{1}{2}}), \quad x_{nk} = \frac{k - np}{\sqrt{npq}}.$$

Výraz $o(n^{-\frac{1}{2}})$ vyjadřuje malý zbytek, který při vzrůstajícím n konverguje k nule rychleji než $n^{-\frac{1}{2}}$. Je tedy zanedbatelný vzhledem k uvedenému členu.

Jest tudíž

$$P \left[\frac{S_{2n}}{2n} = \frac{1}{2} \right] = \binom{2n}{n} 2^{-2n} = \frac{1}{\sqrt{\pi n}} + o(n^{-\frac{1}{2}})$$

pro $S_{2n} \sim Bi(n, \frac{1}{2})$ a k tomuto odhadu lze opět použít Stirlingův vzorec.

V kontextu příkladu 8 tedy platí: Pravděpodobnost toho, že při n hodech mincí obdržíme výsledek 1 (líc) přesně v polovině případů, je řádově malá jako $n^{-\frac{1}{2}}$. Čebyševova nerovnost však v tomto případě vypovídá, že ať je okolí $\frac{1}{2}$ jakkoliv malé, pak relativní četnost $\frac{S_n}{n}$ se nalézá v tomto okolí s pravděpodobností, která je řádově velká jako $1 - n^{-1}$. Tato dvě sdělení mohou přispět ke správnému pochopení zákona velkých čísel.

Věta 6 (Moivreova-Laplaceova věta integrální). Je-li $X_n \sim Bi(n, p)$ pro $p \in (0, 1)$, pak

$$\lim_{n \rightarrow \infty} P \left[a \leq \frac{X_n - np}{\sqrt{npq}} \leq b \right] = P[a \leq Y \leq b]$$

stejněměrně pro $-\infty \leq a \leq b \leq \infty$, kde $Y \sim N(0, 1)$.

Moivreovy-Laplaceovy věty (1801) patří do historie matematiky. Z poloviny minulého století pochází následující vynikající zpřesnění věty integrální:

Věta 7 (Nerovnost Berry-Essénova). *Při okolnostech a značení integrální věty 6 platí*

$$\left| P \left[a \leq \frac{X_n - np}{\sqrt{npq}} \leq b \right] - P[a \leq Y \leq b] \right| \leq 1,6 \frac{p^2 + q^2}{\sqrt{npq}}$$

pro libovolnou volbu $-\infty < a \leq b < \infty$.

Řád chyby $n^{-\frac{1}{2}}$ nelze zlepšit, aproximace binomického rozdělení Gaussovým jsou tím přesnější, čím je pravděpodobnost p blíží $\frac{1}{2}$ (funkce $\frac{p^2 + (1-p)^2}{\sqrt{p(1-p)}}$ má minimum rovno jedné pro $p = \frac{1}{2}$). Pro intervaly typu $(-\infty, b)$ platí Berry-Essénova nerovnost s konstantou 0,8.

Poučení 2. Je-li $X_n \sim Bi(n, p)$ a $n \rightarrow \infty$, aproximujeme $\frac{X_n - np}{\sqrt{npq}} \sim N(0, 1)$ s přesností, která je dána Berry-Essénovou nerovností.

Takto poučení můžeme zkrátit univerzální interval koncentrace pravděpodobnosti (25). Je-li $X \sim Bi(n, p)$ pro n velké, pak

$$\begin{aligned} P[\mu - 1,96\sigma \leq X \leq \mu + 1,96\sigma] &= P \left[\left| \frac{X - np}{\sqrt{npq}} \right| \leq 1,96 \right] \\ &\doteq \frac{1}{\sqrt{2\pi}} \int_{-1,96}^{1,96} e^{-\frac{1}{2}t^2} dt \doteq 0,95, \end{aligned}$$

kde jsme použili obvyklé značení $\mu = \mathbf{E}X = np$ a $\sigma^2 = \text{var}X = npq$.

2. Podmiňování a nezávislost

Začneme opět matematickým modelem.

Definice 7. $(\Omega, \mathcal{F}, P)$ buď pravděpodobnostní prostor, F a G náhodné jevy v $\mathcal{F}$ a $P(G) > 0$. Číslo

$$P(F|G) = \frac{P(F \cap G)}{P(G)}$$

se nazývá podmíněná pravděpodobnost náhodného jevu F při podmínce G (také čteme – pravděpodobnost F podmíněna G).

Definice by měla vyjadřovat, jaký vliv může mít informace, že výsledek pokusu má vlastnost G ($\omega \in G$), na nové posouzení pravděpodobnosti náhodného jevu F . V KPP nad Ω má podmíněná pravděpodobnost tvar

$$P(F|G) = \frac{|F \cap G|}{|G|}, \quad (26)$$

což je nepodmíněná pravděpodobnost jevu, že výsledek pokusu má vlastnost $F \cap G$ v modelu KPP nad $\Omega = G$. Vyzkoušejme, zda definice splňuje naše očekávání.

Vraťme se k příkladům 1 a 2 z první části. Při výběru s vracením by informace o tom, že v prvním tahu byla tažena bílá koule (B_1), měla být zcela irelevantní pro posouzení pravděpodobnosti toho, že i v druhém tahu bude tažena bílá koule (B_2), tj. mělo by platit $P(B_2|B_1) = P(B_2)$. Je tomu tak, neboť

$$\begin{aligned} P(B_2|B_1) &= \frac{P(B_1 \cap B_2)}{P(B_1)} = \frac{|B_1 \cap B_2|}{|B_1|} = \frac{b^2}{b(a+b)} = \frac{b}{a+b} \\ &= P(B_2). \end{aligned} \quad (27)$$

Při výběru bez vracení je informace B_1 podstatná. Tah bílé koule v prvním tahu připravil pro druhý tah osudí s novým barevným složením (a černých koulí, $b-1$ bílých). Mělo by tedy být $P(B_2|B_1) = \frac{b-1}{a+b-1}$. Je tomu tak:

$$P(B_2|B_1) = \frac{|B_1 \cap B_2|}{|B_1|} = \frac{b(b-1)}{b(a+b-1)} = \frac{b-1}{a+b-1}. \quad (28)$$

Vrátíme se ještě k příkladu 3 z předchozí části a určíme podmíněnou pravděpodobnost

$$P[K_2 = k_2 | K_1 = k_1]$$

jevu, že v druhé přihrádce bude k_2 částic, bylo-li již zjištěno, že v první přihrádce je k_1 částic:

$$\begin{aligned} P[K_2 = k_2 | K_1 = k_1] &= \frac{|[K_1 = k_1, K_2 = k_2]|}{|[K_1 = k_1]|} \\ &= \frac{\binom{n}{k_1} \binom{n-k_1}{k_2} (r-2)^{n-(k_1+k_2)}}{\binom{n}{k_1} (r-1)^{n-k_1}} \\ &= \frac{\binom{n-k_1}{k_2} ((r-1)-1)^{(n-k_1)-k_2}}{(r-1)^{n-k_1}}, \end{aligned} \quad (29)$$

což je nepodmíněná pravděpodobnost toho, že v modelu s $r - 1$ přihrádkami a $n - k_1$ částicemi bude v druhé přihrádce k_2 částic: Oznamí-li pozorovatel, že v první přihrádce zjistil k_1 částic, bude podmíněná pravděpodobnost toho, že v druhé přihrádce je k_2 částic počítána jako nepodmíněná pravděpodobnost ve smyslu vzorce (26). Podmíněné pravděpodobnosti umožňují modelování složitějších dvoustupňových experimentů pomocí následujícího jednoduchého vzorce pro úplnou pravděpodobnost. Buď $\Omega = \bigcup_{k=0}^n F_k$ disjunktní rozklad prostoru Ω (tedy $F_i \cap F_j = \emptyset$ pro $i \neq j$) a $P(F_k) > 0$ pro $0 \leq k \leq n$, pak

$$P(B) = \sum_{k=0}^n P(F_k)P(B|F_k) \quad \text{pro} \quad B \in \mathcal{F}. \quad (30)$$

Příklad 10. Deseti bílými či černými koulemi je osudí naplněno tak, že bylo desetkrát hozeno symetrickou mincí; padl-li rub (líc), byla do osudí vložena koule bílá (černá). Z takto náhodně naplněného osudí postupně vybíráme n koulí, taženou koulí do osudí vracíme. Jaká je pravděpodobnost $P(B_n)$ jevu, že všechny tažené koule jsou bílé?

Jde skutečně o dvoustupňový náhodný experiment se složitou strukturou kombinatorického prostoru, který by jej modeloval. Jednodušeji můžeme vstupní informace interpretovat následovně. O barevném složení osudí činíme hypotézy $F_0, F_1, \dots, F_{10}$, kde F_k označuje osudí s k bílými koulemi. Podle příkladu 8 je $P(F_k) = \binom{10}{k} 2^{-10}$. Abychom mohli použít vzorec (30), potřebujeme modelovat podmíněné pravděpodobnosti $P(B_n|F_k)$. Přirozeným modelem je nepodmíněná pravděpodobnost vytažení n bílých koulí s vracením z osudí, kde se nachází k bílých a $10 - k$ koulí. Podle příkladu 9 tedy je $P(B_n|F_k) = \left(\frac{k}{10}\right)^n$. Celkově dostáváme

$$\begin{aligned} P(B_n) &= \sum_{k=0}^{10} \binom{10}{k} 2^{-10} \left(\frac{k}{10}\right)^n = 2^{-10} \sum_{k=0}^{10} \binom{10}{k} \left(1 - \frac{k}{10}\right)^n \\ &\leq 2^{-10} \sum_{k=0}^{10} \binom{10}{k} e^{-\frac{kn}{10}} = 2^{-10} (1 + e^{-\frac{n}{10}})^{10}. \end{aligned} \quad (31)$$

Značný význam má následující zdánlivě primitivní inverze. Nechť $P(B) > 0$ a $P(F) > 0$, pak

$$P(F|B) = \frac{P(F)P(B|F)}{P(B)} = \frac{P(F)P(B|F)}{\sum_{k=0}^n P(F_k)P(B|F_k)} \quad (32)$$

kteřá se nazývá Bayesův vzorec a umožňuje řešit úlohy následujícího typu.

Příklad 11. Uvažte situaci z příkladu 10. Z osudí byly taženy výhradně bílé koule. Jaká je pravděpodobnost toho, že osudí neobsahovalo žádnou kouli černou?

Máme počítat podmíněnou pravděpodobnost $P(F_{10}|B_n)$. Podle vztahů (28) a (31) je

$$P(F_{10}|B_n) = \frac{2^{-10}1}{2^{-10} \sum_{k=0}^{10} \binom{10}{k} \left(\frac{k}{10}\right)^n} \geq \frac{1}{(1 + e^{-\frac{n}{10}})^{10}}$$

a zjišťujeme, jak jsme očekávali, že $\lim_{n \rightarrow \infty} P(F_{10}|B_n) = 1$, speciálně $P(F_{10}|B_{50}) = 0,9504$ nebo $P(F_{10}|B_{100}) = 0,9997$.

V některých úlohách je třeba opatrně interpretovat vstupní údaje jako absolutní, respektive podmíněné pravděpodobnosti.

Příklad 12. Tenista má prvé podání úspěšné s pravděpodobností 0,6, druhé s pravděpodobností 0,8. S jakou pravděpodobností p se hráč dopustí dvojchyby? (řešení: $p = 0,08$.)

Podrobné řešení této úlohy je obsahem příkladu 2.2 v [1], další příklady tohoto typu jsou 2.3, 2.4, 2.5.

Kdybychom chtěli prohlásit dvě vlastnosti výsledku náhodného pokusu F a G za nezávislé, jistě bychom ověřovali rovnosti $P(F|G) = P(F)$ a $P(G|F) = P(G)$, a tedy ekvivalentně, rovnost $P(F \cap G) = P(F)P(G)$ (pokud $P(F) > 0$ a $P(G) > 0$).

Definice 8. Náhodné jevy F a G jsou nezávislé, když platí $P(F \cap G) = P(F)P(G)$. Náhodné veličiny X a Y jsou nezávislé, když rovnost

$$P[X = x, Y = y] = P[X = x]P[Y = y] \quad (33)$$

platí pro $(x, y) \in \mathbb{R}^2$, tj. když každá dvojice $[X = x]$, $[Y = y]$ je dvojicí nezávislých náhodných jevů.

Uvažme nezávislé jevy F a G a počítejme podle pravidel (3) a (4) z části 1. Dostáváme

$$\begin{aligned} P(F^c \cap G) &= P(G - F \cap G) = P(G) - P(F)P(G) \\ &= (1 - P(F))P(G) = P(F^c)P(G). \end{aligned} \quad (34)$$

Podobně snadno nahlédneme, že jestliže (F, G) je dvojice nezávislých jevů, pak i všechny dvojice (F^c, G) , (F, G^c) a (F^c, G^c) jsou dvojicemi nezávislých jevů, takže zjišťujeme, že nezávislost vykazuje určitou stabilitu.

Odsud plyne, že dvě náhodné veličiny $X \sim Alt(p_1)$ a $Y \sim Alt(p_2)$ jsou nezávislé právě tehdy, když

$$P[X = 1, Y = 1] = P[X = 1]P[Y = 1]. \quad (35)$$

Vyzkoušejme, zda definice nezávislosti splňuje naše očekávání. Uvažme náhodné jevy B_1 a B_2 (tah bílé koule v prvním a druhém tahu) z příkladů 1 a 2. Vráť-li se tažená koule do osudí, je jeho barevné složení pro druhý tah stejné jako pro tah první. Jevy B_1 a B_2 by měly být nezávislé a je tomu tak, protože $P(B_2|B_1) = P(B_2)$ podle (2). Nevrací-li se tažená koule, má osudí před druhým tahem jiné barevné složení určené výsledkem tahu prvního. Jevy B_1 a B_2 by nezávislé intuitivně být neměly a opravdu se snadno přesvědčíme, že nejsou, protože

$$P(B_2|B_1) = \frac{b-1}{a+b-1} \neq \frac{b}{a+b} = P(B_2)$$

podle (3).

Uvažme ještě Maxwellův-Boltzmannův model z příkladu 3 v první části, a to s n částicemi a $r = 2$ přihrádkami. Náhodné veličiny K_1 a K_2 , které označují počty částic v první a druhé přihrádce, by neměly být nezávislé, protože $K_1 + K_2 = n$ (lineární závislost). Je tomu tak, protože $P[K_2 = k_2|K_1 = k_1] = 1$, je-li $k_2 = n - k_1$; pro jiná k_2 je tato pravděpodobnost nulová.

Velmi důležitou charakteristikou vztahu mezi náhodnými veličinami X a Y je jejich kovariance.

Definice 9. X a Y buďte náhodné veličiny. Číslo

$$\begin{aligned} \text{cov}(X, Y) &= \mathbf{E}(X - \mathbf{E}X)(Y - \mathbf{E}Y) \\ &= \sum_{(x,y)} (x - \mathbf{E}X)(y - \mathbf{E}Y)P[X = x, Y = y] \\ &= \mathbf{E}(XY) - \mathbf{E}X\mathbf{E}Y \end{aligned} \quad (36)$$

se nazývá kovariance X a Y .

Poznamenejme, že první rovnost je definice, druhá je důsledkem lemmatu 1 (pod větou 1) a třetí rovnost plyne roznásobením a výpočtem podle (16).

Věta 8. Jsou-li X a Y nezávislé NV, pak

$$\mathbf{E}(XY) = \mathbf{E}X\mathbf{E}Y, \text{ cov}(X, Y) = 0 \text{ a } \text{var}(X + Y) = \text{var}X + \text{var}Y. \quad (37)$$

Podle již zmíněného lemmatu je

$$\begin{aligned}\mathbf{E}(XY) &= \sum_{(x,y)} xy \mathbf{P}[X = x, Y = y] \\ &= \sum_{(x,y)} x \mathbf{P}[X = x] \cdot y \mathbf{P}[Y = y] \\ &= \sum_x x \mathbf{P}[X = x] \sum_y y \mathbf{P}[Y = y] = \mathbf{E}X \mathbf{E}Y.\end{aligned}\tag{38}$$

Jelikož $\text{cov}(X, Y) = \mathbf{E}(XY) - \mathbf{E}X \mathbf{E}Y$ podle (36), plyne z nezávislosti také, že $\text{cov}(X, Y) = 0$. Snadno spočítáme, že

$$\text{var}(X + Y) = \mathbf{E}(X + Y - \mathbf{E}(X + Y))^2 = \text{var}X + \text{var}Y + 2 \text{cov}(X, Y)$$

a druhá rovnost implikuje třetí.

Je-li $\text{cov}(X, Y) = 0$ říkáme, že X a Y jsou nekorelované NV.

Příklad 13. Přesvědčte se, že náhodné veličiny X a Y s alternativním rozdělením jsou nezávislé právě tehdy, když jsou nekorelované.

Skutečně, buď $X \sim \text{Alt}(p_1)$, $Y \sim \text{Alt}(p_2)$ a $\text{cov}(X, Y) = 0$. Pak

$$\begin{aligned}\mathbf{P}[X = 1, Y = 1] &= \sum_{(x,y)} xy \mathbf{P}[X = x, Y = y] = \mathbf{E}(XY) = \mathbf{E}X \mathbf{E}Y \\ &= p_1 p_2 = \mathbf{P}[X = 1] \mathbf{P}[Y = 1]\end{aligned}$$

a veličiny X a Y jsou nezávislé podle (35).

Obecně je však nezávislost silnější požadavek než nekorelovanost.

Příklad 14. Nechť X a Y jsou dvě náhodné veličiny takové, že

$$\begin{aligned}\mathbf{P}[X = 1, Y = 1] &= \mathbf{P}[X = 1, Y = -1] = \mathbf{P}[X = -1, Y = -1] \\ &= \mathbf{P}[X = -1, Y = 1] = \mathbf{P}[X = 0, Y = 0] = \frac{1}{5}.\end{aligned}$$

Ukažte, že veličiny X a Y jsou nekorelované, ale jsou závislé.

Určíme rozdělení NV X a Y :

$$\begin{aligned}\mathbf{P}[X = 1] &= \mathbf{P}[X = 1, Y = 1] + \mathbf{P}[X = 1, Y = -1] = \frac{2}{5}, \\ \mathbf{P}[X = -1] &= \frac{2}{5}, \\ \mathbf{P}[X = 0] &= \frac{1}{5}.\end{aligned}$$

Symetricky $P[Y = 1] = P[Y = -1] = \frac{2}{5}$ a $P[Y = 0] = \frac{1}{5}$. Veličiny X a Y nejsou nezávislé, protože

$$\frac{1}{5} = P[X = 0, Y = 0] \neq P[X = 0]P[Y = 0] = \frac{1}{25}.$$

Veličiny X a Y jsou nekorelované, protože ze symetrie jest $\mathbf{E}X = \mathbf{E}Y = 0$ a $\text{cov}(X, Y) = \mathbf{E}(XY)$ je rovna

$$1 \cdot 1 \frac{2}{5} + 1(-1) \frac{2}{5} + (-1)(-1) \frac{2}{5} + (-1)1 \frac{2}{5} + 0 \cdot 0 \frac{1}{5} = 0.$$

Příklad 15. $K_1, K_2, \dots, K_r$ buďte počty částic v přihrádkách $1, 2, \dots, r$ (příklad 3). Již víme, že $K_j \sim Bi(n, \frac{1}{r})$, kde n je celkový počet částic. Tedy jest $\mathbf{E}K_j = \frac{n}{r}$ a $\text{var}K_j = nr(1 - \frac{1}{r})$. Pokuste se vypočítat, že $\text{cov}(K_i, K_j) = -\frac{n}{r^2}$ pro $i \neq j$.

Obecnější výpočet naleznete v příkladu 20.

Poznámka 2. Důležitou mírou závislosti NV X a Y je jejich korelační koeficient

$$\rho(X, Y) = \frac{\text{cov}(X, Y)}{\sqrt{\text{var}X} \sqrt{\text{var}Y}}.$$

Jest $|\rho_{X,Y}| \leq 1$, víme, že $\rho(X, Y) = 0$ pro nezávislé NV X a Y a lze dokázat (viz [1], věta 7.5), že $|\rho(X, Y)| = 1$ platí právě tehdy, když existují konstanty a, b, c takové, že $P[aX + bY = c] = 1$.

Užitečné je následující rozšíření pojmu nezávislosti.

Definice 10. Náhodné jevy $F_1, F_2, \dots, F_n$ jsou nezávislé, když

$$P(F_{k_1} \cap F_{k_2} \cap \dots \cap F_{k_r}) = P(F_{k_1})P(F_{k_2}) \dots P(F_{k_r}) \quad (39)$$

platí pro každou volbu $1 \leq k_1 < k_2 < \dots < k_r \leq n$.

Důležité je vědět, že požadavky (39) nelze redukovat.

Příklad 16. Vraťme se k příkladu 7 a uvažme prostor, který jsme nazvali generátor náhodných čísel v $[0, 1]$. Dokažte, že náhodné jevy

$$F_1 = \left[0, \frac{1}{2}\right], \quad F_2 = \left[\frac{1}{4}, \frac{3}{4}\right] \quad \text{a} \quad F_3 = \left[0, \frac{1}{4}\right] \cup \left[\frac{1}{2}, \frac{3}{4}\right]$$

jsou nezávislé po dvou, ale nikoliv nezávislé ve smyslu předchozí definice.

Volíme-li $F_1 = F_2 = [0, \frac{1}{2}]$ a $F_3 = [\frac{1}{2}, \frac{1}{2}]$, vidíme, že (39) nelze redukovat na požadavek $P(\bigcap_1^n F_k) = P(F_1)P(F_2) \cdots P(F_n)$.

Indukcí snadno rozšíříme platnost stability nezávislosti.

Budte $F_1, F_2, \dots, F_n$ nezávislé jevy, pak také $G_1, G_2, \dots, G_n$ jsou nezávislé jevy při každé volbě $G_k = F_k$ nebo $G_k = F_k^c$.

Definice 11. Náhodné veličiny $X_1, X_2, \dots, X_n$ jsou nezávislé, když

$$P[X_{k_1} = x_1, \dots, X_{k_r} = x_r] = P[X_{k_1} = x_1] \cdots P[X_{k_r} = x_r] \quad (40)$$

platí při každé volbě $1 \leq k_1 < k_2 < \dots < k_r \leq n$ a $(x_1, x_2, \dots, x_r) \in \mathbb{R}^r$.

S nezávislostí více než dvou náhodných veličin jsme se již setkali.

Příklad 17. Uvažte posloupnost NV $X_1, X_2, \dots$ definovaných na prostoru, který jsme nazvali generátor náhodných čísel v $[0, 1]$ a $X_n(\omega)$ je n-tý člen dvojkového rozkladu čísla $\omega \in [0, 1]$. Rovnost (39) říká, že $X_n \sim Alt(\frac{1}{2})$. Dokažte, že pro každé $n \in \mathbb{N}$ jsou náhodné veličiny $X_1, X_2, \dots, X_n$ nezávislé.

Snadno například ověříme, že platí

$$\begin{aligned} P[X_1 = 1, X_2 = 1, \dots, X_n = 1] &= P\left(\omega \in \left[\sum_{k=1}^n 2^{-k}, 1\right]\right) \\ &= 1 - \sum_{k=1}^n 2^{-k} = 2^{-n} = P[X_1 = 1]P[X_2 = 1] \cdots P[X_n = 1]. \end{aligned}$$

Příklad 18. Vraťme se k příkladu 8. Buď S_n počet hodů s výsledkem 1 (líc mince). Zřejmě je $S_n = \sum_{k=1}^n X_k$, kde X_k je nula nebo jedna tak, že $X_k = 1$ právě tehdy, když k-tý hod zaznamenal výsledek 1. Necht' jsou všechny výsledky náhodného pokusu, zřejmě tvořené posloupnostmi nul a jedniček délky n , stejně pravděpodobné. Ukažte, že pak platí $X_k \sim Alt(\frac{1}{2})$ a NV $X_1, X_2, \dots, X_n$ jsou nezávislé.

Ověříme nezávislost pro dvě a tři NV, dále lze postupovat analogicky. Pro $k < l < j$ platí

$$\begin{aligned} P[X_k = 1] &= \frac{2^{n-1}}{2^n} = \frac{1}{2}, \\ P[X_k = 1, X_l = 1] &= \frac{2^{n-2}}{2^n} = \frac{1}{4} \\ &= P[X_k = 1]P[X_l = 1], \\ P[X_k = 1, X_l = 1, X_j = 1] &= \frac{2^{n-3}}{2^n} = \frac{1}{8} \\ &= P[X_k = 1]P[X_l = 1]P[X_j = 1]. \end{aligned}$$

V první části jsme ukázali, že $S_n = \sum_{k=1}^n X_k$ je NV s binomickým rozdělením $Bi(n, \frac{1}{2})$. Toto je obecnější zákonitost. Buďte $X \sim Bi(n, p)$ a $Y \sim Bi(m, p)$ dvě nezávislé NV. Pak

$$\begin{aligned}
 P[X + Y = k] &= \sum_{l=0}^n P[X = l, Y = k - l] \\
 &= \sum_{l=0}^n P[X = l]P[Y = k - l] \\
 &= \sum_{l=0}^n \binom{n}{l} \binom{m}{k-l} p^l (1-p)^{n-l} p^{k-l} (1-p)^{m-(k-l)} \\
 &= p^k (1-p)^{n+m-k} \sum_{l=0}^n \binom{n}{l} \binom{m}{k-l} \\
 &= \binom{n+m}{k} p^k (1-p)^{n+m-k}
 \end{aligned}$$

platí pro $0 \leq k \leq n+m$. Použili jsme konvenci $\binom{n}{k} = 0$ pro $k > n$. Dokázali jsme, že platí implikace

$$X \sim Bi(n, p), Y \sim Bi(m, p) \text{ nezávislé} \Rightarrow X + Y \sim Bi(n+m, p) \quad (41)$$

a dokonce i tvrzení

Věta 9. $X_1, X_2, \dots, X_n$ buďte nezávislé NV takové, že každá z nich má alternativní rozdělení $Alt(p)$. Jejich součet $S_n = \sum_k X_k$ má pak binomické rozdělení $Bi(n, p)$.

Důkaz se provede indukcí. Implikace (41) zajišťuje platnost tvrzení pro $n = 2$, neboť $Alt(p) = Bi(1, p)$ a tak víme, že $X_1 + X_2 \sim Bi(2, p)$. Necht' tvrzení platí pro $n - 1$, dokážeme platnost pro n indukcí. Pro n napíšeme $X_1 + X_2 + \dots + X_n = (X_1 + \dots + X_{n-1}) + X_n$, kde $X_1 + \dots + X_{n-1} \sim Bi(n-1, p)$ podle indukčního předpokladu a $X_n \sim Bi(1, p)$; snadno ověříme, že tyto dvě náhodné veličiny jsou nezávislé a implikace (41) nás přivádí k závěru, že $X_1 + X_2 + \dots + X_n \sim Bi(n, p)$.

Poučení 3. Náhodná veličina X s alternativním rozdělením $Alt(p)$ je nepochybně vhodným modelem dichotomického pokusu s výsledkem úspěch (1), jehož pravděpodobnost je p , a neúspěch (0) s pravděpodobností $q = 1 - p$. Ukázali jsme, že počet úspěchů S_n při n nezávislých opakováních takového

pokusu má binomické rozdělení $Bi(n, p)$. Podle Čebyševovy nerovnosti I (věta 3) je pro $\varepsilon > 0$

$$P \left[\left| \frac{S_n}{n} - p \right| < \varepsilon \right] \geq 1 - \frac{pq}{\varepsilon^2 n} \geq 1 - \frac{1}{4\varepsilon^2 n}, \quad (42)$$

protože $\mathbf{E} \frac{S_n}{n} = \frac{np}{n} = p$ a $\text{var} \frac{S_n}{n} = \frac{1}{n^2} npq = \frac{pq}{n}$, jak víme z minula. Správná, třeba neznámá, hodnota pravděpodobnosti úspěchu p je v ε -okolí relativní četnosti úspěchů $\frac{S_n}{n}$ s pravděpodobností, která je nejméně $1 - \frac{1}{4\varepsilon^2 n}$.

Poučení 1 a 2 tedy říkají: při velkém počtu n nezávislých opakování dichotomického pokusu aproximujeme rozdělení počtu úspěchů S_n rozdělením Poissonovým $Po(np)$, je-li pravděpodobnost p malá. V opačném případě aproximujeme rozdělení normovaného počtu úspěchů $\frac{S_n - np}{\sqrt{npq}}$ rozdělením normálním $N(0, 1)$ s přesností, kterou udává Berry-Essénova nerovnost.

Ve větě 8 jsme ukázali, že pro nezávislé NV X a Y je $\text{var}(X + Y) = \text{var}X + \text{var}Y$. Indukcí, podobně jako ve větě 9, dostáváme

Věta 10. $X_1, X_2, \dots, X_n$ buďte nezávislé NV. Pak

$$\text{var} \left(\sum_{k=1}^n X_k \right) = \sum_{k=1}^n \text{var} X_k. \quad (43)$$

Poznamenejme, že jsme znovu dokázali rovnosti (24). Jestliže platí $X \sim Bi(n, p)$, můžeme předpokládat, že $X = \sum_{k=1}^n X_k$, kde X_k jsou nezávislé NV. Dále platí

$$\mathbf{E}X = \sum_{k=1}^n \mathbf{E}X_k = np, \quad \text{var}X = \sum_{k=1}^n \text{var}X_k = npq.$$

Můžeme také rozšířit působnost nerovnosti (42) následujícím způsobem.

Věta 11 (Čebyševova nerovnost II). $X_1, X_2, \dots, X_n$ buďte nezávislé NV se stejným rozdělením pravděpodobností. Označíme $\mathbf{E}X_k = \mu$, $\text{var}X_k = \sigma^2$. Pak pro každé $\varepsilon > 0$ platí nerovnost

$$P \left[\left| \frac{1}{n} \sum_{k=1}^n X_k - \mu \right| \geq \varepsilon \right] \leq \frac{\sigma^2}{\varepsilon^2 n}. \quad (44)$$

Poznamenejme, že předpoklad o stejném rozdělení veličin X_k , tj. předpoklad $P[X_1 = x] = P[X_2 = x] = \dots = P[X_n = x]$ pro $x \in \mathbb{R}$, triviálně implikuje, že

$$\mathbf{E}X_1 = \mathbf{E}X_2 = \dots = \mathbf{E}X_n = \mu, \text{ var}X_1 = \text{var}X_2 = \dots = \text{var}X_n = \sigma^2.$$

Důkaz věty 11 je snadný. Podle vět 1 a 2 a (43) spočteme ve větě 9 $\mathbf{E}\frac{X}{n} = \mu$ a $\text{var}\frac{X}{n} = \frac{n\sigma^2}{n^2} = \frac{\sigma^2}{n}$ a aplikujeme prvou Čebyševovu nerovnost, abychom obdrželi (44).

Poznamenejme, že máme-li k dispozici celou posloupnost $X_1, X_2, \dots$ nezávislých NV se stejným rozdělením pravděpodobností a označíme-li $\mathbf{E}X_k = \mu$, pak

$$\lim_{n \rightarrow \infty} P \left[\left| \frac{1}{n} \sum_{k=1}^n X_k - \mu \right| < \varepsilon \right] = 1, \quad \varepsilon > 0. \quad (45)$$

Jakkoliv chaotické je chování náhodné posloupnosti $X_1, X_2, \dots$, její postupné aritmetické průměry „konvergují“ ke společné střední hodnotě μ ve smyslu (45). Příklad 17 takovou posloupnost konstruuje. Je-li $\omega = \sum_{k=1}^{\infty} \frac{X_k(\omega)}{2^k}$ dvojkový rozvoj $\omega \in [0, 1]$, pak $X_k \sim \text{Alt}(\frac{1}{2})$ a $X_1, X_2, \dots$ je posloupnost nezávislých NV. Zjistili jsme, že aritmetické průměry $\frac{1}{n} \sum_{k=1}^n X_k$ „konvergují“ k $\frac{1}{2}$ ve smyslu (45). Do nerovnosti (44) vstupuje vektor náhodných veličin $(X_1, X_2, \dots, X_n)$, které jsou nezávislé a mají stejná rozdělení pravděpodobností. Obecněji definujeme pojem náhodného vektoru.

Definice 12. $(\Omega, \mathcal{F}, P)$ buď pravděpodobnostní prostor, $X_1, X_2, \dots, X_n$ zde definované NV. Zobrazení $\mathbf{X} = (X_1, X_2, \dots, X_n)$ definované na Ω s hodnotami v $\mathbb{R}^n$ se nazývá n -rozměrný náhodný vektor. Funkce

$$p(x_1, x_2, \dots, x_n) = P[X_1 = x_1, X_2 = x_2, \dots, X_n = x_n],$$

pro argument $(x_1, x_2, \dots, x_n) \in \mathbb{R}^n$ se nazývá rozdělení pravděpodobností náhodného vektoru $\mathbf{X}$.

Je zřejmé, že funkce $p(x_1, x_2, \dots, x_n)$ může být rozdělením některého náhodného vektoru pouze tehdy, když

$$\begin{aligned} p(x_1, x_2, \dots, x_n) &= 0 \text{ až na konečně mnoho } (x_1, x_2, \dots, x_n), \\ p(x_1, x_2, \dots, x_n) &\in [0, 1] \\ \sum_{(x_1, \dots, x_n) \in \mathbb{R}^n} p(x_1, x_2, \dots, x_n) &= 1. \end{aligned} \quad (46)$$

Z pravděpodobnostního hlediska náhodný vektor není pouze souborem náhodných veličin. Toto ukazuje příklad 14 a také příklad následující.

Příklad 19. $\mathbf{X} = (X_1, X_2)$ buď dvourozměrný náhodný vektor takový, že

$$P[\mathbf{X} = (1, 1)] = P[\mathbf{X} = (0, 1)] = P[\mathbf{X} = (1, 0)] = P[\mathbf{X} = (0, 0)] = \frac{1}{4}.$$

$\mathbf{Y} = (Y_1, Y_2)$ buď dvourozměrný náhodný vektor takový, že

$$P[\mathbf{Y} = (1, 1)] = P[\mathbf{Y} = (0, 0)] = \frac{1}{2}.$$

Přesvědčte se, že vektory $\mathbf{X}$ a $\mathbf{Y}$ nemají stejná rozdělení pravděpodobností, i když jejich souřadnice stejně rozdělené jsou. Souřadnice X_1 a X_2 jsou nezávislé, souřadnice Y_1 a Y_2 nezávislé nejsou.

Poučení je, že rozdělení náhodného vektoru není jednoznačně určeno tím, že zadáme rozdělení jednotlivých souřadnic. Poučení dále je, že rozdělení libovolného náhodného vektoru $(X_1, X_2, \dots, X_n)$, řekněme $p(x_1, x_2, \dots, x_n)$, je jednoznačně určeno tím, že zadáme rozdělení každé z NV $X_1, X_2, \dots, X_n$ a přidáme požadavek, aby tyto NV byly nezávislé. Je tomu tak proto, že v tomto případě je

$$\begin{aligned} p(x_1, x_2, \dots, x_n) &= P[X_1 = x_1, X_2 = x_2, \dots, X_n = x_n] \\ &= P[X_1 = x_1]P[X_2 = x_2] \cdots P[X_n = x_n]. \end{aligned}$$

Netriviální příklad náhodného vektoru je vektor s multinomickým rozdělením. Označme

$$S_{nr} = \left\{ (k_1, k_2, \dots, k_r), \quad 0 \leq k_j \leq n, \quad \sum_{j=1}^r k_j = n, \quad k_j \in \mathbb{Z} \right\}.$$

Definice 13. Náhodný vektor $\mathbf{X} = (X_1, X_2, \dots, X_r)$ s hodnotami v množině S_{nr} má multinomické rozdělení $MN(n, r, p_1, \dots, p_r)$, kde $0 \leq p_j \leq 1$ a $\sum_{j=1}^r p_j = 1$, jestliže

$$\begin{aligned} P[X_1 = k_1, X_2 = k_2, \dots, X_r = k_r] &= \\ &= \binom{n}{k_1} \binom{n-k_1}{k_2} \cdots \binom{n-k_1-\cdots-k_{r-1}}{k_r} p_1^{k_1} p_2^{k_2} \cdots p_r^{k_r} \\ &= \frac{n!}{k_1! k_2! \cdots k_r!} p_1^{k_1} p_2^{k_2} \cdots p_r^{k_r}. \quad (47) \end{aligned}$$

Vrátíme-li se k Maxwellovu-Boltzmannovu modelu a uvažíme náhodné veličiny $K_1, K_2, \dots, K_r$, které udávají počty částic $k_1, k_2, \dots, k_r$ v přihrádkách $1, 2, \dots, r$, vypočítáme, že pro $(k_1, k_2, \dots, k_r) \in S_{nr}$ je

$$P[K_1 = k_1, K_2 = k_2, \dots, K_r = k_r] = \frac{\binom{n}{k_1} \binom{n-k_1}{k_2} \binom{n-k_1-k_2}{k_3} \dots 1}{r^n}.$$

Platí tedy $(K_1, K_2, \dots, K_r) \sim \text{MN}(n, r, \frac{1}{r}, \dots, \frac{1}{r})$ a zkoumaný náhodný vektor má multinomické rozdělení.

Poznamenejme, že součet pravděpodobností (47) je jedna (tak jak má být), protože podle multinomické věty platí

$$\sum_{(k_1, \dots, k_r) \in S_{nr}} \binom{n}{k_1} \binom{n-k_1}{k_2} \dots \binom{k_r}{k_r} p_1^{k_1} p_2^{k_2} \dots p_r^{k_r} = (p_1 + p_2 + \dots + p_r)^n \quad (48)$$

pro libovolnou volbu $n \in \mathbb{N}$, $r \in \mathbb{N}$ a $p_j \in \mathbb{R}$, $j = 1, 2, \dots, r$.

Příklad 20. Uvažte vektor $(X_1, X_2, \dots, X_r) \sim \text{MN}(n, r, p_1, \dots, p_r)$ a vypočtěte kovarianci $\text{cov}(X_i, X_j)$.

Pro $0 \leq k_1 \leq n$ označme $\sum^* = \sum_{(k_2, \dots, k_r) \in S_{n-k_1, r-1}}$ a počítejme

$$\begin{aligned} P[X_1 = k_1] &= \sum^* P[X_1 = k_1, X_2 = k_2, \dots, X_r = k_r] \\ &= \sum^* \binom{n}{k_1} \binom{n-k_1}{k_2} \dots \binom{n-\sum_{j=2}^{r-1} k_j}{k_r} p_1^{k_1} p_2^{k_2} \dots p_r^{k_r} \\ &= \binom{n}{k_1} p_1^{k_1} (1-p_1)^{n-k_1} \end{aligned}$$

podle (48), kde volíme $n = n - k_1$, $r = r - 1$ a $p_2, p_3, \dots, p_r$. Jest tedy

$$X_j \sim \text{Bi}(n, p_j), \mathbf{E}X_j = np_j \text{ a } \text{cov}(X_j, X_j) = \text{var}X_j = np_j(1-p_j).$$

Obdobně, pro $0 \leq k_1 + k_2 \leq n$, vypočítáme:

$$P[X_1 = k_1, X_2 = k_2] = \binom{n}{k_1} \binom{n-k_1}{k_2} p_1^{k_1} p_2^{k_2} (1-p_1-p_2)^{n-k_1-k_2}$$

a jsme schopni určit $\text{cov}(X_1, X_2)$. Nejprve vypočteme $\mathbf{E}(X_1 X_2)$ podle postupu uvedeného v definici 6.

$$\begin{aligned}
\mathbf{E}(X_1 X_2) &= \sum_{0 \leq k_1 + k_2 \leq n} k_1 k_2 P[X_1 = x_1, X_2 = x_2] \\
&= \sum_{2 \leq k_1 + k_2 \leq n} \frac{n! p_1^{k_1} p_2^{k_2} (1 - p_1 - p_2)^{n - k_1 - k_2}}{(k_1 - 1)!(k_2 - 1)!(n - k_1 - k_2)!} \\
&= n(n - 1)p_1 p_2 \sum_{0 \leq l_1 + l_2 \leq n - 2} \frac{(n - 2)! p_1^{l_1} p_2^{l_2} (1 - p_1 - p_2)^{n - 2 - l_1 - l_2}}{l_1! l_2! (n - 2 - l_1 - l_2)!} \\
&= n(n - 1)p_1 p_2 (p_1 + p_2 + 1 - p_1 - p_2)^{n - 2} = n(n - 1)p_1 p_2,
\end{aligned}$$

opět podle (48), protože

$$\frac{(n - 2)!}{l_1! l_2! (n - 2 - l_1 - l_2)!} = \binom{n - 2}{l_1} \binom{n - 2 - l_1}{l_2}.$$

Odsud

$$\text{cov}(X_1, X_2) = \mathbf{E}(X_1 X_2) - \mathbf{E}X_1 \mathbf{E}X_2 = n(n - 1)p_1 p_2 - np_1 np_2 = -np_1 p_2.$$

Obecněji pro $i \neq j$ dostáváme $\text{cov}(X_i, X_j) = -np_i p_j$.

Pokud se vám nepodařilo vyřešit příklad 15, dostáváme řešení nyní. Pro složky K_i a K_j náhodného vektoru s multinomickým rozdělením platí, že pro $i \neq j$ je $\text{cov}(K_i, K_j) = -\frac{n}{r^2}$.

Poznámka k literatuře

Základy počtu pravděpodobnosti lze studovat z nepřehledného množství knih. V příspěvku jsme používali odkaz na skripta Zvára, Štěpán (2003) určená pro studenty učitelství na MFF UK. Mezi další možné zdroje poučení lze zařadit skripta Dupač, Hušková (1999) určená pro studenty všech matematických oborů na MFF UK, starší skripta Likeš, Machek (1981), ale i první části klasických učebnic Feller (1967), Gněděnka (1969) či Rényiho (1972). Mnoho zajímavých příkladů vhodných i pro studenty středních škol obsahuje kniha Anděl (2000). Přehled nejstarší historie pravděpodobnosti inspirované hazardními hrami lze najít v Mačák (1997).

Literatura

- [1] Zvára K. a Štěpán J. (2003). *Pravděpodobnost a matematická statistika*, MATFYZPRESS, Praha.
- [2] Dupač V. a Hušková M. (1999). *Pravděpodobnost a matematická statistika*, Karolinum, Praha.
- [3] Likeš J. a Machek J. (1981). *Počty pravděpodobnosti*, SNTL, Praha.
- [4] Feller W. (1967). *Introduction to Probability Theory and Its Applications I*, Wiley, Chichester. (Existuje dostupnější ruský překlad)
- [5] Gněděnko B.V. (1969). *Kurs teorii věrojatnostěj*, Mir, Moskva. (Anglicky vyšlo 1976)
- [6] Rényi A. (1972). *Teorie pravděpodobnosti*, Academia, Praha.
- [7] Anděl J. (2000). *Matematika náhody*, MATFYZPRESS, Praha.
- [8] Mačák K. (1997). *Počátky počtu pravděpodobnosti*, Prometheus, Praha.

GEOMETRICKÁ PRAVDĚPODOBNOST

Ivan Saxl

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

ivan.saxl@mff.cuni.cz

1. Úvod

Geometrická pravděpodobnost a prostorová statistika se vyvíjely ještě pomaleji a obtížněji, než pravděpodobnost a statistika dějů, jež lze jednoduše numericky popsat. Tato skutečnost je o to paradoxnější, že k získání poznatků o náhodných číselných dějích bylo nutné začít házet kostky či hrát karty, zatím co organický i anorganický svět, který nás obklopuje, je přeplněn pravděpodobnostními jevy, jejichž realizacemi jsou nejrůznější geometricky popsatelné objekty - izolované, jako jsme my sami, živočichové a rostliny se svými plody, nebo naopak úzce navzájem související s cílem vyplnit část prostoru, jako jsou krystaly hornin, buňky tkání, lesní a luční porosty aj. První poznatky o rozdělení odchylek byly získány při měření polohy hvězd, avšak až u Galilea se setkáváme s myšlenkou, že odchylky od nejčastější hodnoty jsou většinou malé a jen zřídka velké (opět ve vztahu k astronomickým měřením). Přitom k podobnému poznatku bychom dospěli porovnáním plodů libovolného stromu, jejichž sběrem se po tisíciletí zabývá počet lidí, ve srovnání s nímž je množství astronomů měřících polohy hvězd zcela zanedbatelné. Možnost aplikovat Gaussův zákon odchylek a posléze i jiné podobné zákony na živé bytosti vyvolala dnes stěží pochopitelné nadšení sociologů (L.A.J. Quetelet) i biologů (F. Galton, W.F.R. Weldon) v 19. století a první polovině století dvacátého. Zdá se tedy, že představa nejistoty v zákonitě určených rozmezích je obtížně přijatelná lidskému duchu přesto, že nejelementárnější zraková zkušenost mu přináší bezpočet poznatků o tom, jak všechno kolem nás je rozmanitě skoro stejné v rozměrech a tvarech celků i jejich složek. Předmětem geometrické pravděpodobnosti a statistiky je zkoumání prostorových jevů: jevem se rozumí skutečnost, že v jisté oblasti prostoru se vyskytuje něco, co můžeme

Klíčová slova: Geometrická pravděpodobnost, úloha o jehle, Croftonův vztah, Hadwigerův teorém, valuace

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839 a grantu GAČR č. 201/03/0946.

popsat geometrickými charakteristikami typu objem, plošný obsah, lineární rozměr a tvar, přičemž se tyto vlastnosti při eukleidovském pohybu zachovávají. K nim přistupuje ještě orientace, která se zachovává pouze při pohybu translačním. Existuje zde jistá analogie s časovými procesy; u nich se v jistém čase něco stane (rozpadne se atom, spadne meteorit, dojde k nehodě), zatímco u prostorových jevů, jež také často nazýváme procesy, se v nějaké oblasti prostoru vyskytuje objekt či soubor objektů.

Jev může být unikátní, např. když se pokoušíme poznat patologickou změnu konkrétního tělesného orgánu z jeho geometrie (to se dobře daří např. u ledvin a v zažívacím ústrojí) nebo analyzovat porušení strojních součástí (nalézt trhlinu, předpovědět její šíření) při nebezpečí havárie nebo po ní. V tomto případě je náhodná (avšak cíleně volená a opakovaná s ohledem na předběžné znalosti) volba místa zkoumání, tj. odběr vzorku tkáně či materiálu.

Druhý typ problémů se objevuje při sledování kolektivních jevů, kdy vedle velikosti objektů (nečistoty, příměsi či dutinky v materiálu, stromy v lese) se zajímáme také o jejich rozmístění, které vypovídá něco o jejich interakci (např. výskyt a rozmístění dřevin určitého typu ve vztahu ke složení půdy a vodním tokům, zalidnění krajiny a jeho vztah k obslužnosti). Extrémně koordinované jevy – teselace, mozaiky – pokrývající bez překrývání části prostoru nalézáme nejen v materiálech a tkáních, ale i u produktů společenské organizace (parcely, správní oblasti, státní celky aj.). Při řešení těchto úloh můžeme často za náhodný považovat zkoumaný jev a místo zkoumání je do značné míry libovolné.

Odhad vlastností prostorových objektů a vztahů resp. dějů mezi nimi provádíme obecně *geometrickým výběrem*. Ten spočívá v analýze jejich *projekcí* do lineárních podprostorů (roviny, přímký) a *průniků* s nimi či s vybranými množinami – *sondami* – známých vlastností, které se předepsaným způsobem pohybují. Obsahy a pravděpodobnosti těchto průniků a obsahy projekcí jsou hlavními objekty teorie i aplikací geometrické pravděpodobnosti. Významnou vlastností většiny úloh je, že dimenze prostoru, v němž děj či zkoumání probíhá, dimenze objektů a sond či projekcí jsou primárními charakteristikami řešeného problému a že tedy jsou základními *proměnnými veličinami*. Jinými slovy, geometrická pravděpodobnost je multidimenzionální disciplína, v níž dimenze interagujících komponent *explicitně* vstupují do všech obecných vztahů.

V prvních dvou kapitolách jsou probírány jednoduché a často historické úlohy s pevně zadanými dimenzemi prostředí i interagujících komponent. Úlohy této části jsou vesměs zcela názorné a mohou být probírány v rámci středoškolské výuky, pro níž jistě budou zajímavým obohacením. Ukazuje

se v nich, že i geometrická pravděpodobnost zpočátku souvisela s hrami, avšak velmi brzy byla využívána k řešení praktických problémů, především v krystalografii a při popisu kovových materiálů. Po překonání zásadních problémů, souvisejících s možnou mnohoznačností úloh, probíhal intenzivní rozvoj teorie i aplikací po celé minulé století, především však v jeho posledních čtyřiceti letech. Velmi rychle byly nalezeny zcela zásadní aplikace v biologii, analýze a interpretaci obrazů, v jejich počítačovém zpracování a úpravách, v půdním průzkumu a v ekologii. Její pronikání do učebních osnov je však neobyčejně pomalé a obecné povědomí o jejích postupech a zákonitostech prakticky neexistuje.

Jednoduché postupy úvodních tří kapitol jsou v následujícím textu zobecněny pro eukleidovské prostory libovolné dimenze d . Přitom si není třeba představovat $d > 3$, ale všimnout si, co mají společného i rozdílného prostory dimenzí $d = 1, 2, 3$, s nimiž běžně přicházíme do styku. Například Eulerova-Poincarého charakteristika je zcela unikátní geometrickou vlastností objektů nezávislou na dimenzi prostoru, v nichž je zkoumáme, nemění se při transformacích objektů (typu deformace bez roztržení) a podporující geometrickou představivost. Tento přístup sice patří již do učební látky vysokoškolské, některé poznatky uvedené v dalších kapitolách, především izoperimetrické nerovnosti, Cauchyho vztah v rovině i trojrozměrném prostoru a metody odhadu plošného obsahu nepravidelných ploch a délek složitých křivkových systémů (pro svou jednoduchost použitelné i v běžném životě) by neměly pro svou obecnou užitečnost chybět ani ve středoškolské geometrii.

2. Elementární geometrické události a jejich pravděpodobnost

Typické úlohy geometrické pravděpodobnosti se týkají náhodných interakcí geometrických objektů, tj. pravděpodobností situací $A \uparrow B \equiv A \cap B \neq \emptyset$ (A zasahuje B), dále $A \subset B$, a konečně $A \not\cap B \equiv A \cap B = \emptyset$ (A míjí B). Pravděpodobnost definujeme podobně jako při numerických situacích jako poměr příznivých případů ke všem možným případům, ty však musíme měřit, nikoliv počítat. Uvažujme úsečku $A = [0, 1]$ rozdělenou na dvě části $A' = [0, 0.1)$, $A'' = [0.1, 1]$ a uvažujme pravděpodobnosti, že bod zasahující A zasáhne A' nebo A'' . Pokud bychom ke stanovení obou pravděpodobností použili běžný způsob určení poměru počtu příznivých případů ke všem možným případům, dospěli bychom k závěru, že oba případy jsou stejně pravděpodobné s pravděpodobností rovnou jedné. Mohutnosti množin A', A'' i A jsou totiž stejné – rovny mohutnosti kontinua $2^{\aleph_0}$. Intuitivně však soudíme, že náhodný bod zasahující A zasáhne v průměru devětkrát častěji A'' než A' ,

protože A'' je devětkrát delší než A' . Pravidlo, že pro účely geometrické pravděpodobnosti porovnáváme nikoliv počty situací, ale jejich míry, je tedy zcela přirozené. V tomto případě měly všechny srovnávané množiny situací stejnou dimenzi $d = 1$ jako prostor $\mathbb{R}^1$, v němž děj probíhal, a mohli jsme porovnat jejich délky. V obecné úloze tohoto typu porovnáváme *Lebesgueovy míry* $\nu_d(A)$, které jsou v $\mathbb{R}^d$ nulové pro všechny množiny, jejichž dimenze $\dim(A) < d$. Pravděpodobnost, že náhodný bod zasahující čtverec A zasáhne také jeho vybranou tětivu b , okraj nebo dokonce vrchol, je tedy 0. Aby míra všech uvažovaných situací byla konečná, musejí být geometrické pravděpodobnosti vhodně podmíněné. Např. pravděpodobnost, že bod x roviny F zasahuje obrazec A libovolné konečné plochy v této rovině ležící, je zřejmě 0, neboť při jejím výpočtu porovnáme obsah $\nu_d(A)$ s obsahem celé roviny. Typický případ bude tedy např. hledání $P(A \uparrow B | A \subset C)$ pro $B \subset C$, $\nu_d(C) < \infty$. Uvedme několik jednoduchých příkladů:

Příklad 1. Určete pravděpodobnost $P(d)$, že bod x zasahující d -rozměrnou krychli A , zasahuje také d -rozměrnou kouli B krychli vepsanou a určete limitu této pravděpodobnosti pro sudé $d \rightarrow \infty$.

Úlohu stačí řešit pro jednotkovou kouli (poloměr $r = 1$); její objem je $\kappa_d = \pi^{d/2} / \Gamma(\frac{d}{2} + 1)$, kde $\Gamma(\bullet)$ je Γ -funkce; pro celé n je $\Gamma(n + 1) = n!$ a $\kappa_d = 1, 2, \pi, \frac{4}{3}\pi, \dots$ pro $d = 0, 1, 2, 3, \dots$. Potom $P(d) = P(x \in B | x \in A) = \kappa_d / 2^d$ a s použitím Stirlingovy formule $d! \approx \sqrt{2\pi d} d^{d+\frac{1}{2}} e^{-d}$ pro sudé d je

$$P_{AB}(d) = \left(\frac{\pi e}{2d}\right)^{d/2} \frac{1}{\sqrt{\pi d}}, \quad \text{tj. } \lim_{d \rightarrow \infty} P_{AB}(d) = 0.$$

Úlohu lze zobecnit na případ libovolné množiny $B \subset A$ v $\mathbb{R}^d$, v němž určujeme $P(x \in B | x \in A) = \nu_d(B) / \nu_d(A)$. Lze ji chápat jako hledání informace o množině B , které provedeme následujícím způsobem: zvolíme libovolný bod, např. počátek souřadného systému x_0 a pohybujeme jím v prostoru. Gruppu eukleidovských translací označíme $\mathcal{G}_t$, její prvky g_t . Ze všech možných translací $g_t x_0$ si budeme všimnout pouze těch, pro něž $g_t x_0 \in A$; označíme je např. $g_{t,A}$, mírou jejich množiny $\mathcal{G}_{t,A}$ je $\nu_d(A)$. Budeme zaznamenávat případy, kdy $g_{t,A} x_0 \in B$; ty označíme $g_{t,AB}$, mírou jejich množiny $\mathcal{G}_{t,AB}$ je $\nu_d(B)$. Poté vytvoříme rovnoměrně náhodný výběr z translací $g_{t,A}$, což znamená, že každá z těchto translací má stejnou šanci, že bude vybrána. Pravděpodobnost, že vybraná translace $g_{t,A}$ je zároveň také $g_{t,AB}$ je $\nu_d(B) / \nu_d(A)$. Přitom množiny A, B jsou zde pevné, náhodný je výběr translací, resp. jím odpovídajících bodů $g_{t,A} x_0$. Pokud jich vybereme N a z nich n bude $g_{t,AB}$, pak n/N bude nestranným odhadem (značíme hranatými závorkami $[\]$) $\nu_d(B) / \nu_d(A)$ a píšeme $[\nu_d(B) / \nu_d(A)] = n/N$. Známe-li navíc obsah $\nu_d(A)$, získáme nestranný

odhad $\nu_d(B)$. Přitom množina B nemusí být souvislá; může být sjednocením několika izolovaných částí, takže metodu můžeme použít např. k odhadu poměrného zastoupení zalesněné plochy na vybrané zmapované oblasti libovolného tvaru nebo k odhadu zalesněné plochy v oblasti známé velikosti. Pohybující se bod x_0 obvykle nazýváme *sondou*.

Lze si představit i obrácenou úlohu. Zvolíme pevnou množinu A , např. dostatečně velký čtverec ve srovnání s listy, a v něm N náhodně nebo systematicky vybraných (např. tvořících čtvercovou bodovou mřížku) pevných bodů. Poté na tento základ budeme jeden po druhém házet listy B_i , $i = 1, \dots, j$, nějakého stromu a zaznamenávat počty listy překrytých bodů n_i . Výraz

$$Q = \frac{\nu_2(A)}{jN} \sum_{i=1}^j n_i$$

je nestranným odhadem středního obsahu proměřených listů. Jednotlivé listy můžeme chápat jako realizace náhodného procesu „list B našeho stromu S “ a Q je nestranný (pokud jsme listy rovnoměrně náhodně vybrali) odhad očekávaného obsahu $\mathbf{E}\nu_2(B)$. V této úloze je tedy náhodným studovaný objekt; ten sám provádí výběr z pevného systému bodů, který obvykle nazýváme *testovacím systémem*, v tomto případě *bodovým*. Popsaný postup se nazývá *bodová metoda* a široce se používá v metalografii i v biologii k odhadu ploch zobrazených strukturních prvků resp. jejich řezů.

Podobně lze postupovat při pravděpodobnostních úlohách na zakřivených plochách, nejčastěji to bývá na sféře. Lebesgueovu míru však musíme nahradit k -rozměrnou *mírou Hausdorffovou* h_k , která pro $k = d$ je shodná s mírou Lebesgueovou a pro ostatní hodnoty k udává např. plošný obsah hranice množiny či délku křivky. Následující příklad je jednoduchou demonstrací takové úlohy.

Příklad 2. Uvažujme sféru $\mathbb{S}^2$ poloměru r a jeden její vrchlík V o výšce v . Jaká je pravděpodobnost, že rovnoměrně náhodně zvolený bod $x \in \mathbb{S}^2$ bude ležet ve vrchlíku V ?

Zřejmě je $P(x \in V | x \in \mathbb{S}^2) = h_2(V)/h_2(\mathbb{S}^2) = \frac{1}{2}p$, kde $p = v/r$ (plocha vrchlíku je $2\pi r v$).

Sonda ovšem nemusí být bodová, jak ukazují následující dvě úlohy.

Příklad 3. Ve středu čtverce A o straně a je kruh B o poloměru $R \leq a/2$. Určete pravděpodobnost, že kruh C o poloměru r zasahující čtverec A , zasáhne také kruh B .

Jako referenční body kruhů zvolíme jejich středy x_B (ten je zároveň středem čtverce) a x_C . Aby kruh C zasáhl čtverec A , musí být vzdálenost $d(x_C, A) = \inf_{a \in A} D(x_C, a) \leq r$; $D(x_C, a)$ je eukleidovská vzdálenost

$\|x_C - a\|$. Všechny polohy středu x_C , pro něž je tato podmínka splněna vymezují tzv. *vnější rovnoběžnou množinu* A_r . Její obsah je $\nu_2(A_r) = a^2 + 4ar + \pi r^2$ (obecně pro konvexní množinu X v $\mathbb{R}^2$ s hranicí délky $h_1(\partial X)$ je $\nu_2(X_r) = \nu_2(X) + rh_1(\partial X) + \pi r^2$, což je speciální případ *Steinerovy věty* - viz níže rovnice (8)). Aby kruh C zasahoval kruh B , musí jeho střed ležet ve vnější rovnoběžné množině B_r kruhu B , jejíž obsah je $\pi(r + R)^2$. Protože $B \subset A$, je také $B_r \subset A_r$. Takže

$$P(C \uparrow B | C \uparrow A) = \frac{\nu_2(B_r)}{\nu_2(A_r)} = \frac{\pi(Q^2 + 2qQ + q^2)}{1 + 4q + \pi q^2}, \text{ kde } q = \frac{r}{a}, Q = \frac{R}{a}.$$

Příklad 4. Jaká je pravděpodobnost, že kruh C poloměru r se středem $x_C \in A$ protíná hranici ∂A čtverce A o straně $a \geq 2r$?

Aby platilo $C \uparrow \partial A$, nesmí být x_C vzdáleno od hranice čtverce více než r a tedy musí ležet ve vnitřním okrajovém pásu čtverce o šířce r . Množina poloh středů x_C takových, že $C \subset A$, je *vnitřní rovnoběžná množina* A_{-r} , takže zmíněný okrajový pás je rozdílem $A - A_{-r}$. Potom pravděpodobnost $P(C \uparrow \partial A | x_C \in A) = 4q(1 - q)$, kde $q = r/a$.

Příklad 5 (Buffonova úloha o čtverci). Oblíbená renesanční hra na čtverce spočívala v házení mince C o poloměru r na podlahu vydlážděnou čtvercovými dlaždicemi. Bodově byly hodnoceny případy, kdy obvod mince protnul systém spár S mezi dlaždicemi ve dvou, třech a čtyřech bodech. Jaký je poměr pravděpodobností $P(\nu_0(\partial C \cap S) = 0)$: $P(\nu_0(\partial C \cap S) = 2)$: $P(\nu_0(\partial C \cap S) = 3)$: $P(\nu_0(\partial C \cap S) = 4)$?

Stačí sledovat mince se středem x_C ležící v jedné dlaždici A o straně a . Hledané pravděpodobnosti jsou potom v poměru obsahů $\nu_2(A_{-r}) : \nu_2(A - A_{-r} - A') : \nu_2(\partial A') : \nu_2(A')$, kde A' je čtverec o straně $2r$ (sjednocení čtyř rohových čtverečků o stranách délky r rovnoběžných se stranami dlaždice A a majících s ní společný jeden vrchol). $\nu_0(\partial C \cap S) = 3$, když se střed mince pohybuje po vnitřních hranách rohových čtverečků - její okraj pak jednu hranu dlaždice protíná ve dvou bodech a druhé se dotýká, Lebesgueova míra těchto situací $\nu_2(\partial A') = 0$. S využitím výsledku předchozí úlohy je potom hledaný poměr $(1 - 2q)^2 : 4q(1 - 2q) : 0 : 4q^2$ (např. pro $q = 0.1$ je to 0.64: 0.32: 0: 0.04). Lze si snadno představit, jakým zdrojem sporů mezi hráči byla nepravděpodobná, leč od jiných těžko odlišitelná situace tří průsečíků.

Jednoduché řešení všech těchto úloh je umožněno tím, že pohybující se či náhodně položené sondy byly d -koule, jejich hranice S^{d-1} nebo body, takže jejich interakce s ostatními množinami byla jednoznačně určena polohou jediného vztažného bodu. V případě sondy obecnějšího tvaru (úsečka, přímka, kružnice v $\mathbb{R}^2$ či v $\mathbb{R}^3$) je však nezbytné zavést také pravidla pro měření

orientací sondy, jež již nejsou tak jednoznačná. Na tuto skutečnost upozornil v knize *Calcul des probabilités* (1888) francouzský matematik J. L. F. Bertrand (1822 - 1900) několika úlohami, které se nazývají Bertrandovy paradoxy. Nejznámější z nich je následující.

Uvažujme kružnici o poloměru r protnutou „náhodnou“ přímkou a hledejme pravděpodobnost, že tětiva je delší než strana rovnostranného trojúhelníku do kružnice vepsaného. Bertrand navrhl tři řešení:

a) Aby tětiva byla delší, musí její střed ležet v kruhu o poloměru $r/2$. Jeho plocha je rovna čtvrtině plochy celé kružnice a tedy hledaná pravděpodobnost je $1/4$.

b) Tětivu určíme pomocí jejích koncových bodů: první bod zvolíme ve vrcholu trojúhelníku A ; zjistíme, že tětiva delší než jeho strana musí protínat jeho stranu BC a opouštět kružnici v bodě oblouku $\widehat{BC}$. Jeho délka je právě jedna třetina délky kružnice a tedy hledaná pravděpodobnost je $1/3$.

c) Tětiva delší než strana trojúhelníku musí protínat koncentrickou kružnici o poloměru $r/2$ (kružnice vepsaná uvažovanému trojúhelníku) a tedy také průměr kružnice kolmý ke svému směru. Mírou delších tětiv je pak poměr $(r/2) : r$ a hledaná pravděpodobnost je $1/2$.

Pro nalezení pravidel, která by odstranila tyto a další mnohoznačnosti, hledejme nejprve míru svazku T přímek v rovině se středovým úhlem ϕ . Přímký tvořící svazek mají jeden společný bod (zde střed kružnice o poloměru r) a vyplňují nějaký středový úhel ϕ . Protože každá přímka protíná kružnici ve dvou bodech, mohli bychom jako vhodnou míru zvolit poloviční délku obou svazkem vytknutých oblouků kružnice, tj. ϕr . Pak by ovšem týž soubor přímek měl míru tím větší, čím větší by byl poloměr kružnice, na které bychom měřili. Vhodné pravidlo pro měření orientací je měřit vždy na kružnici jednotkového poloměru, tedy na jednotkové sféře $\mathbb{S}^1$; odpovídající míru označíme $\gamma_1^2 = \frac{1}{2}h_1(T \cap \mathbb{S}^1)$, kde horní index označuje dimenzi prostoru a spodní dimenzi přímek. Mírou izotropního svazku přímek (všechny možné orientace přímek jsou rovnoměrně zastoupeny) pak bude π . Pokud budou přímky orientované, bude jich dvojnásobný počet a tedy mírou orientovaných směrů bude 2π . Podobně je v $\mathbb{R}^3$ mírou směrů vymezených trsem T přímek $\gamma_1^3(T) = \frac{1}{2}h_2(T \cap \mathbb{S}^2)$ (vyplňují dvojité neomezený kužel; termín *trs* se běžně používá pouze pro dimenzi $d = 3$) procházejících středem jednotkové sféry $\mathbb{S}^2$ polovina obsahu $h_2(T \cap \mathbb{S}^2)$ plochy tímto trsem na sféře vytknuté, resp. celý tento obsah v případě orientovaných směrů. Pro izotropní svazky T v $\mathbb{R}^d$ budou tyto míry rovny $\gamma_1^d(T) = \mathcal{O}_{d-1} = d\kappa_d$ pro soubor orientovaných směrů, resp. $\frac{1}{2}\mathcal{O}_{d-1}$ pro směry neorientované; zde $\mathcal{O}_{d-1} = h_{d-1}(\mathbb{S}^{d-1})$ je obsah jednotkové sféry $\mathbb{S}^{d-1}$, tj. hranice d -koule o poloměru 1. Připomeňme, že pro $d = 1, 2, \dots$ je $\mathcal{O}_{d-1} = d\kappa_d = 2, 2\pi, 4\pi, \dots$; κ_d bylo zavedeno v příkladu 1.

Další pravidlo je třeba navrhnout pro charakterizaci svazku T rovnoběžných přímk v $\mathbb{R}^d$. Terminologie není zcela jednotná; v $\mathbb{R}^2$ nazýváme svazkem množinu přímk různých orientací majících jeden společný bod, analogickou množinu v $\mathbb{R}^3$ nazýváme trs, v $\mathbb{R}^d$ kužel. Zde dáváme přednost společnému termínu svazek, a to i pro „pás“ či „válec“ tvořený přímkami rovnoběžnými, které mají společný úběžný bod. Začneme s nejjednodušším případem svazku rovnoběžných přímk v rovině; označme L přímku, která je se svazkem rovnoběžná a prochází počátkem $\mathcal{O}$ zvoleného souřadného systému. Všechny rovnoběžné přímky bez ohledu na svou polohu v prostoru jsou tak reprezentovány jednou přímkou - lineárním podprostorem jednotkové dimenze. V zásadě bychom svazek mohli měřit délkou úsečky, kterou svazek vytkne na libovolné přímce vzhledem ke svazku pevně orientované. Přirozené je opět zvolit tuto míru minimální, tj. jako délku $h_1(T \cap L_\perp)$ vytknutou svazkem na přímce $L_\perp$ procházející počátkem a mající směr normály svazku; používá se pro ni názorný termín *šířka svazku*. Můžeme ji také chápat jako délkou ortogonální projekce $(T|L_\perp)$ do komplementárního (k L) lineárního podprostoru $L_\perp$. Zásadní pravidlo je, že všechny směry i všechny polohy přímk daného směru jsou měřeny shodně, bez ohledu na jiné efekty s jejich polohami či orientacemi spojené, jakými je např. početnost či velikost jejich průniků s hodnocenými objekty,

Bezprostředním zobecněním do $\mathbb{R}^3$ je pravidlo, že mírou (válcového) svazku T přímk rovnoběžných s L je obsah $h_2(T|L_\perp)$ jejich ortogonální projekce $(T|L_\perp)$ do komplementárního lineárního podprostoru (ekvivalentně do normálové roviny). Podobně svazek rovnoběžných rovin T_2 je reprezentován obsahem průniku $T_2 \cap L_{2\perp}$ svazku s jeho normálou či ekvivalentně obsahem průmětu $T_2|L_{2\perp}$ svazku do jeho normálového lineárního podprostoru. Zavedená pravidla demonstrují následující příklady.

Příklad 6. Dva svazky rovnoběžných přímk protínají celou jednotkovou úsečku A a svírají s ní úhly $0 \leq \theta_1, \theta_2 \leq \pi/2$. Určete šířky svazků w_1, w_2 .

Normály obou svazků svírají s úsečkou úhly doplňkové k θ_1 a θ_2 , takže $w_1 = \sin \theta_1, w_2 = \sin \theta_2$.

Příklad 7. Úsečku A délky a zasahuje svazek rovnoběžných přímk T šířky w kolmých k A . Úsečka $B \subset A$ má délku b . Jaká je pravděpodobnost $P(T \uparrow B|T \uparrow A)$, že svazek T zasahuje také úsečku B ?

Jedná se vlastně o jednorozměrnou úlohu typu: Jaká je pravděpodobnost, že úsečka $W = T \cap A$ délky w , zasahující úsečku A délky a , zasáhne také úsečku $B \subset A$ délky b ? Protože úsečky jsou vlastně jednorozměrné koule, řešíme analogicky, jako příklad 3. Aby $W \uparrow A$, musí střed úsečky W ležet ve vnější rovnoběžné množině $A_{w/2}$, aby $W \uparrow B$, musí střed úsečky W le-

žet ve vnější rovnoběžné množině $B_{w/2}$ a hledaná pravděpodobnost je podíl $\nu_1(B_{w/2})/\nu_1(A_{w/2}) = \frac{b+w}{a+w}$.

Příklad 8. Úsečku A délky a protíná v jejím středu úsečku B délky $b \leq a$ svírající s ní úhel θ . Úsečku A dále protíná přímka F k ní kolmá. Jaká je pravděpodobnost $P = P(F \uparrow B | F \uparrow A)$, že tato přímka protíná také úsečku B ?

Opět se jedná o jednorozměrný případ. Přímka F „vidí“ úsečku B jako její projekci délky $b \cos \theta$ do své normály, která má v tomto případě směr shodný s A , takže $P = \frac{b \cos \theta}{a}$.

Všimněme si nyní Bertrandových řešení. Začneme s případem c) a označíme délku strany vepsaného trojúhelníku t_Δ . Každá tětiva resp. přímka ji generující je charakterizována jedním bodem na průměru kruhu k ní kolmému, tj. nezávisle na své délce. Mírou všech tětiv stejného směru je úsečka vytknutá na přímkce obsahující jejich normálu - průměr kružnice $2r$. Mírou tětiv delších než t_Δ je polovina tohoto průměru. Tento výsledek platí pro všechny směry tětiv, míry situací jsou určeny ve shodě s navrženými pravidly a hledaná pravděpodobnost je $P = \frac{1}{2}$.

Jednodušší a intuitivnější je úvaha následující: zvolme jako pevnou přímku osu y a posouvajme kružnici o poloměru r po ose x . Přímka pak bude protínat kružnici tehdy, když její střed $c \in [-r, r]$, a tětiva bude delší než t_Δ pro $c \in [-\frac{r}{2}, \frac{r}{2}]$. Odtud opět je hledaná pravděpodobnost $P = \frac{1}{2}$.

V případě a) je každá tětiva charakterizována svým středem, který pro všechny tětivy různých směrů a stejné délky t leží na kružnici poloměru r_t , $0 \leq r_t \leq r$. Pro tětivy délek $t \geq t_\Delta$ je $r_t \leq r/2$. Soubory tětiv různých délek, a tedy i přímky je generující, jsou tedy měřeny různě: mírou těch nejdelších je 0 (mají středový bod společný, totožný se středem kruhu), mírou těch nejkratších (tečny s tětivami nulových délek) je celý obvod kruhu $2\pi r$. To je v rozporu s přijatými pravidly, tětiv kratších než t_Δ jsou pak 3/4, delších pouze 1/4, a řešení tedy není správné. Podobně je tomu v případě b). Obecné pravidlo zní, že polohy i směry sond i bodů, přímek či rovin zasahujících objekt měříme podle toho, jakou vyplňují v prostoru oblast a nezávisle na jejich konkrétním průniku s objektem.

Následující úloha je obecně považována za první závažnou úlohu z geometrické pravděpodobnosti.

Příklad 9 (Buffonova úloha o jehle). Na podlahu tvořenou prkny šířky d oddělenými spárou zanedbatelné šířky házíme (izotropně náhodně) jehlu B délky $\ell < d$. Jaká je pravděpodobnost, že jehla protne spáru?

Podle předcházející úlohy „uvidí“ přímka F představující spáru jehlu B jako její projekci do své normály. Svírá-li jehla se spárou úhel θ , bude délka

této její projekce $\ell_{\perp} = \ell \sin \theta$. Úhel θ je rovnoměrně náhodný na $[0, \pi/2]$ a očekávaná (střední) délka průřezu jehly bude

$$\mathbf{E}\ell_{\perp} = \frac{2\ell}{\pi} \int_0^{\pi/2} \sin \theta d\theta = \frac{2\ell}{\pi},$$

kde $d\theta$ bychom mohli psát také jako $d\gamma_1^2$. Dále se již jedná opět o jedno-rozměrnou úlohu následujícího typu: na přímce je systém bodů A_i o souřadnicích $\pm id, i = 0, 1, 2, \dots$, navzájem vzdálených o d . Jaká je pravděpodobnost, že rovnoměrně náhodně vybraný interval I délky $w = \mathbf{E}\ell - \perp$ zasáhne některý z bodů A_i . Aby k zásahu došlo, musí střed I ležet v intervalu $(-w/2 + id, id + w/2)$. Pravděpodobnost této situace je zřejmě rovna poměru w/d a dosazením za w dostaneme $P = \frac{2\ell}{\pi d}$.

Buffonovy úlohy o čtverci a jehle jsou historicky první úlohy o geometrické pravděpodobnosti a umožňují velmi dalekosáhlá zobecnění, jak ukazuje příklad 10 a následující diskuse.

Příklad 10. Na systém S rovnoběžných spár z příkladu 9 jsme vysypali neznámý počet jehel neznámých různých délek a dostali jsme celkem N průsečíků. Odhadněte celkovou délku L všech jehel.

$N = \frac{2}{\pi d} \sum_i n_i \ell_i = \frac{2}{\pi d} L$, kde n_i jsou počty jehel délek ℓ_i . Odtud odhad $[L] = \pi d N / 2$.

V předcházejícím příkladu jsme určili celkovou délku rozsypaných jehel bez znalosti jejich počtů a délek. Můžeme si tedy představit, že jehly nám s libovolnou přesností aproximují nějakou křivku či celý systém křivek o délce L a jeho délku jsme schopni odhadnout bez měření pouhým počítáním průsečíků se spárami. Výsledek by platil, i kdyby spáry byly od sebe různě vzdálené se střední vzdáleností $\mathbf{E}d$, tj. $[L] = \pi \mathbf{E}d N / 2$.

Uvažovaný systém spár tedy hraje úlohu čarového testovacího systému, který má délkovou intenzitu (střední délku na jednotkové ploše) $L_A^S = 1/d$ nebo obecněji $L_A^S = 1/\mathbf{E}d$. Sledujme nyní „izotropní“ křivku (všechny orientace jejích tečen jsou stejně zastoupeny) délky L na libovolné oblasti obsahu A . Označíme-li $L_A = L/A$ (formálně můžeme opět L_A považovat za délkovou intenzitu), dostaneme dosazením do vztahu pro L

$$L_A = \frac{\pi}{2} N_L,$$

kde jsme $N_L = N/(AL_A^S)$ označili počet průsečíků na jednotkové délce čar testovacího systému L_A^S (počet průsečíků byl N a celková délka testovacího systému je AL_A^S). Získaná rovnice je užitečným výsledkem použití geometrické

pravděpodobnosti. Umožňuje nám velmi snadno odhadovat délku složitých čárových systémů tak, že na zkoumanou plochu obsahu A (např. fotografický snímek čárových složek mikrostruktury či mapu vodních toků) pokládáme soustavu rovnoběžných čar a počítáme průsečíky testovacího systému se studovaným čárovým systémem L . Odhad jeho délky je potom $[L] = \frac{\pi}{2} AN_L$. Předpokládali jsme, že systém L je izotropní; pokud tomu tak není, musíme učinit izotropním systém L_A^S tím, že jej položíme několikrát s různou orientací – buď náhodnou nebo systematicky měněnou.

3. Sylvestrův vztah

Buffonovu úlohu o jehle můžeme vyřešit i jiným a pro teorii geometrické pravděpodobnosti charakteristickým způsobem (Barbier, 1860). Pro jehlu délky ℓ_1 obecně větší než d je počet průsečíků náhodná proměnná X_1 . Její střední hodnota je

$$\mathbf{E}X_1 = \sum_{n \geq 0} np_n,$$

kde p_n je pravděpodobnost, že průsečíků je právě n . Speciálně, jestliže $\ell_1 < d$, pak $\mathbf{E}X_1 = \mathbf{0} \cdot \mathbf{p}_0 + \mathbf{1} \cdot \mathbf{p}_1 = \mathbf{p}_1$ je hledaná pravděpodobnost. Počet průsečíků druhé nezávislé házené jehly délky ℓ_2 bude další nezávislá náhodná proměnná X_2 . Jestliže jehly budou mít jeden společný koncový bod, nebudou X_1 a X_2 již nezávislé, nicméně stále platí

$$\mathbf{E}(X_1 + X_2) = \mathbf{E}X_1 + \mathbf{E}X_2$$

a analogická rovnice bude platit i pro lomenou čáru o délce $\ell_1 + \ell_2 + \dots$ tvořenou libovolným počtem spojených jehel. Označme $f(\ell)$ funkci vyjadřující závislost středních hodnot počtů průsečíků na délkách jehel. Pak platí $f(X_1 + X_2 + \dots) = f(X_1) + f(X_2) + \dots$ a $f(\ell)$ je lineární monotónní rostoucí funkce typu $f(\ell) = r\ell$ pro libovolné $\ell \geq 0$. Lomenou čáru lze považovat za aproximaci křivky C délky ℓ , jejíž počet průsečíků bude náhodná proměnná Y přibližně rovná $X_1 + X_2 + \dots$. Přejdem k limitě $\ell_i \rightarrow 0$ pro všechna i dostaneme $\mathbf{E}(Y) = r\ell$, kde r je konstanta. Její hodnotu určíme např. když za C zvolíme kružnici o průměru d . Počet průsečíků $r\pi d$ pak bude při každé její poloze roven dvěma, takže $r = \frac{2}{\pi d}$ a tedy $\mathbf{E}(Y) = p_1 = \frac{2\ell_1}{\pi d}$ pro $\ell_1 < d$ jako v příkladu 9. Typičnost tohoto postupu spočívá v tom, že nejprve odvodíme obecnou relaci úměrnosti mezi charakteristikami obecného, obvykle konvexního geometrického objektu, a náhodné proměnné jím generované na obecném testovacím systému. Dosazením jednoduchého objektu, obvykle d -rozměrné koule nebo odpovídající sféry pak určíme konstantu úměrnosti. Postup lze demonstrovat na dalším podobném příkladu.

Označme si ϕ_1^2 invariantní míru (viz níže) definovanou na množině všech přímek $\mathcal{F}_1^2$ v $\mathbb{R}^2$. Uvažujme úsečku délky L_1 a necht' $Z_1 = 1$ resp. 0 je její počet průsečíků s náhodně vybranou přímkou $F_1^2 \in \mathcal{F}_1^2$. Integrál

$$\int_{\mathcal{F}_1^2} Z_1 d\phi_1^2$$

je mírou všech přímek, které ji protínají. Jeho hodnota závisí opět pouze na délce L_1 , a podobně jako u Buffonovy úlohy můžeme přejít k lomené čáře $L_1 + L_2 + \dots$ aproximující křivku C délky L_C a zavést funkci $f(L_C) = \int_{\mathcal{F}_1^2} Z_C d\phi_1^2 = rL_C$, kde Z_C je počet průsečíků křivky C s náhodnou přímkou z $\mathcal{F}_1^2$. Jestliže je nyní křivka C hranicí ∂K konvexní množiny K , je Z_C buď 2 nebo 0 (tečné přímky lze zanedbat) a tedy $f(L_C) = \int_{\mathcal{F}_1^2} Z_C d\phi_1^2 = 2\phi_1^2(K) = rL(\partial K)$, kde $\phi_1^2(K)$ je míra všech přímek protínajících K . Mějme nyní druhou konvexní množinu $K' \subset K$. Pro ni můžeme napsat podobnou rovnici a vydělením obou rovnic dostaneme Sylvestrův vztah

$$\frac{\phi_1^2(K')}{\phi_1^2(K)} = P(F_1^2 \uparrow K' | F_1^2 \uparrow K) = \frac{L(\partial K')}{L(\partial K)}.$$

Zcela analogicky bychom mohli místo přímek F_1^2 uvažovat body F_0^2 s mírou ϕ_0^2 , mající zřejmě význam plošného obsahu, takže např. $\phi_0^2(K) = A(K)$. Sylvestrův vztah má potom tvar

$$P(F_0^2 \uparrow K' | F_0^2 \uparrow K) = \frac{\nu_2(K')}{\nu_2(K)},$$

který jsme již intuitivně použili v příkladu 1 a při odhadu plochy listů.

4. Orientace sond

Po počátečních úvahách věnovaných odděleně úlohám na přímce, v rovině a v prostoru, se nyní pokusme řešit úlohy geometrické pravděpodobnosti obecně v $\mathbb{R}^d$. Když dimenze množiny není uvedena explicitně, je rovna d , tj. $\dim(K) = d$ a $\dim(B_r) = r < d$. Pokud je třeba zdůraznit, k prostoru jaké dimenze se množina či funkcionál vztahují, je přidán horní index - např. F_r^d je r -rovina v $\mathbb{R}^d$. Důvodem k obecnému postupu je, že dimenze je pak explicitní proměnnou a že řada vztahů je pochopitelnější ve tvaru pro obecné d , než pro jednotlivé konkrétní dimenze. Bude se jednat o popis interakce množiny A se zvolenou pohybující se resp. náhodně položenou a orientovanou množinou B_r , kterou jsme nazvali sonda. Často je konvexní, může to však být i úsek křivky (např. cykloidy) nebo n -tice bodů, jako byl bodový testovací systém. Především to však bude přímka či rovina, resp. obecně r -rovina F_r^d v $\mathbb{R}^d$ („1-rovina“ je vždy přímka, $(d-1)$ -rovina je *nadrovina*). Pro popis pohybu sondy

B_r^d v ní zvolíme pevný kartézský souřadnicový systém $\Omega = (\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_r)$, popsaný d souřadnicemi jeho jednotkových vektorů, a umístíme jej do zvoleného počátku (referenčního bodu) b_0 . Translace sondy pak budou popsány všemi polohami $g_t b_0$, kde g_t je prvek grupy eukleidovských translací v $\mathbb{R}^d$. Jedné orientaci sondy tedy odpovídá jeden bod prostoru $\mathbb{R}^{dr}$. Jeho souřadnice v $\mathbb{R}^{dr}$ nejsou nezávislé, ale splňují $\binom{r+1}{2}$ relací $\mathbf{e}_i \mathbf{e}_j = \delta_{ij}$, kde $\delta_{ij} = 1$ pro $i = j$ a 0 jinak (Kroneckerovo δ). Všechny možné orientace sondy jsou pak popsány varietou dimenze $s = (rd - r(r+1)/2)$ v $\mathbb{R}^{dr}$, kterou nazýváme *Stiefelova varieta* $\mathcal{S}_r^d$; její s -obsah označíme $h_s(\mathcal{S}_r^d)$. S jejími jednoduchými příklady jsme se již setkali. Variety $\mathcal{S}_1^d$ popisující orientaci úsečky jsou jednotkové sféry $\mathbb{S}^{d-1}$, neboť orientace je popsána jediným vektorem $\mathbf{e}_1$ se souřadnicemi $\{e_{1,1}, \dots, e_{1,d}\}$, splňujícími jedinou relaci $e_{1,1}^2 + \dots + e_{1,d}^2 = 1$. Avšak již orientaci obrazce v $\mathbb{R}^2$ popisuje varieta $\mathcal{S}_2^2$ dimenze 1 v $\mathbb{R}^4$: dvojice vektorů $\mathbf{e}_1, \mathbf{e}_2$ je v $\mathbb{R}^4$ popsána čtyřmi souřadnicemi $\{e_{1,1}, e_{1,2}, e_{2,1}, e_{2,2}\}$ splňujícími tři rovnice $e_{1,1}^2 + e_{1,2}^2 = e_{2,1}^2 + e_{2,2}^2 = 1$, $e_{1,1}e_{2,1} + e_{1,2}e_{2,2} = 0$. Ty určují křivku (dimenze $s = 4 - 2 \cdot 3/2 = 1$) v $\mathbb{R}^4$ a její délka je $h_1(\mathcal{S}_2^2) = 4\pi$, neboť mírou všech orientací vektoru $\mathbf{e}_1$ je 2π a $\mathbf{e}_2$ má právě dvě možnosti odpovídající pravo- a levotočivé soustavě. Podobně mírou orientací tělesa v $\mathbb{R}^3$ je $\mathcal{S}_3^3$ v $\mathbb{R}^9$ s obsahem $h_3(\mathcal{S}_3^3) = 4\pi \times 2\pi \times 2 = \mathcal{O}_2 \mathcal{O}_1 \mathcal{O}_0 = 16\pi^2$ a dimenzí $s = 3$. Obecně je $h_s(\mathcal{S}_r^d) = \mathcal{O}_{d-1} \mathcal{O}_{d-2} \dots \mathcal{O}_{d-r}$. Dané soubory orientací popisují oblasti na těchto varietách; jejich obsahy jsou míry těchto souborů a budeme je značit σ_r^d . Obsah infinitesimální oblasti $d\mathcal{S}_r^d$ značíme $d\sigma_r^d$ místo $h_s(d\mathcal{S}_r^d)$.

Když sondou je r -rovina F_r^d , budeme ve shodě s počátečními úvahami její polohu určovat průsečíkem s komplementárním lineárním podprostorem - $(d-r)$ -rovinou procházející počátkem, kterou budeme značit $L_{r\perp}^d$. Zřejmě je $\dim(L_{r\perp}^d) = (d-r)$, takže budeme psát L_{d-r}^d , pokud vztah k r -rovinám není třeba zdůrazňovat (jako např. při projekcích $K|L_k^d$). Mírou poloh svazku T_r rovnoběžných r -rovin pak bude obsah $h_{d-r}(T_r \cap L_{r\perp}^d)$ odpovídající $(d-r)$ -rozměrné oblasti $T_r \cap L_{r\perp}^d \equiv T_r|L_{r\perp}^d$ komplementu $L_{r\perp}^d$. Chceme-li nyní popsat všechny možné orientace dané r -roviny, musíme si uvědomit, že rotace souřadného systému uvnitř F_r , určené např. změnami prvních r souřadnic každého z vektorů $\mathbf{e}_i$ a popsané tedy Stiefelovou varietou $\mathcal{S}_r^r$, orientaci r -roviny nemění. Všechny možné orientace r -roviny popisuje tedy jiná varieta $\mathcal{G}_r^d$, nazývaná *Grassmanova*; její dimenze g je zřejmě rozdíl dimenzí $\mathcal{S}_r^d$ a $\mathcal{S}_r^r$, tj. $g = r(d-r)$, a její obsah je podíl obsahů $\mathcal{S}_r^d$ a $\mathcal{S}_r^r$, tj.

$$h_g(\mathcal{G}_r^d) = \mathcal{O}_{d-1} \mathcal{O}_{d-2} \dots \mathcal{O}_{d-r} / \mathcal{O}_{r-1} \mathcal{O}_{r-2} \dots \mathcal{O}_0.$$

Pro míru celého souboru orientací jsme již zavedli značení γ_r^d pro $h_g(\mathcal{G}_r^d)$, obsah infinitesimální oblasti $d\mathcal{G}_r^d$ značíme $d\gamma_r^d$.

5. Aditivní funkcionály (valuace)

Obsahem geometrické pravděpodobnosti je sledování interakcí náhodných množinových systémů; je proto třeba zavést charakteristiky, které nám sdělí, zda k události $B_r \uparrow A \equiv B_r \cap A \neq \emptyset$ došlo, a dále jaký je její rozsah, tj. průnik $B_r \cap A$ nějak „změří“. Omezíme se na systém $\mathcal{K}$ všech konvexních množin a dále na systém jejich konečných sjednocení $\mathcal{R}$, který nazýváme *konvexní okruh*. *Aditivním funkcionálem*, resp. *valuací* (modernější termín) nazveme množinovou funkci μ definovanou primárně nad $\mathcal{K}$ a splňující pro každé $A, B \in \mathcal{K}$ podmínku $\mu(A \cup B) = \mu(A) + \mu(B) - \mu(A \cap B)$ a dále $\mu(\emptyset) = 0$. Indukcí pak dostaneme

$$\begin{aligned} \mu(A_1 \cup \dots \cup A_n) = & \sum_i^n \mu(A_i) - \sum_{i < j} \mu(A_i \cap A_j) + \sum_{i < j < k} \mu(A_i \cap A_j \cap A_k) - \\ & - \dots + (-1)^{n-1} \mu(A_1 \cap A_2 \cap \dots \cap A_n). \end{aligned} \quad (1)$$

Jestliže $\mu(gK) = \mu(K)$, kde $g \in \mathcal{G}$ je prvek eukleidovské pohybové grupy (translace, reflexe, rotace), nazýváme μ *pohybově invariantním* či jen *invariantním* aditivním funkcionálem. Aditivní funkcionál nazýváme *spojitým*, když $\mu(A_n) \rightarrow \mu(A)$ pro posloupnost množin A_n s limitou A (tj. Hausdorfova vzdálenost

$$\delta(A_n, A) = \max\left(\sup_{a_n \in A_n} d(a_n, A), \sup_{a \in A} d(a, A_n)\right)$$

množin A_n , A konverguje k nule pro $n \rightarrow \infty$; přitom

$$d(x, A) = \inf_{a \in A} D(x, a)$$

($D(x, a)$ jsme již zavedli jako eukleidovskou vzdálenost mezi body x, a). Funkcionál nazveme *monotónním*, když $\mu(A) \leq \mu(B)$ pro každé $A \subset B$. Konečně jestliže $\mu(\alpha K) = \alpha^k \mu(K)$, je μ *homogenní* stupně k . Příklady invariantních aditivních funkcionálů nad $\mathcal{K}$ jsou obsahy množin a jejich hranic a také Eulerova charakteristika, která má mezi všemi funkcionály tohoto typu výlučné postavení, vyjádřené následující větou:

Věta 1. *Existuje jediná spojitá invariantní valuace μ_0^d nad konvexním okruhem $\mathcal{R}$ taková, že $\mu_0^d(K) = 1$ pro libovolnou neprázdnou konvexní kompaktní množinu $K \in \mathcal{K}$.*

Z její aditivity plyne, že pro soubor množin $A_1, A_2, \dots, A_n$ takových, že $A_i \cap A_j = \emptyset$ pro každou dvojici $i \neq j$, $i, j = 1, \dots, n$, platí, že $\mu_0^d(A_1 \cup A_2 \cup \dots \cup$

$A_n) = n$. Přitom jestliže K leží v nějakém podprostoru (r -rovině) dimenze r prostoru $\mathbb{R}^d$, pak $\mu_0^r(K) = \mu_0^d(K)$, takže můžeme psát $\mu_0(K)$ místo $\mu_0^d(K)$. Jinými slovy, Eulerova charakteristika udává počet izolovaných neprázdných konvexních množin libovolné dimenze v jejich konečném sjednocení $A \in \mathcal{K}$. Bude tedy tím funkcionálem, který popíše, zda k události $B_r \uparrow A$ došlo – pak $\mu_0(B_r \uparrow A) \neq 0$, případně že vzniklo právě $\mu_0(B_r \uparrow A)$ izolovaných komponent tohoto průniku (množinou konvexního okruhu jsou např. tělesa tvaru žebříku nebo činky, jejichž průniky s rovinou či přímkou mají obecně několik izolovaných konvexních komponent).

Požadavek konvexity není nijak striktní. Hodnota Eulerovy charakteristiky se totiž nezmění při homeomorfní transformaci množiny K (tj. při zobrazení, jež je spojitým přechodem mezi vzorem a obrazem, jakým je např. i nehomogenní deformace). Eulerova charakteristika úsečky je 1 a nemění se při jejím zakřivení (žádné dva body však nesmějí splynout). Podobně můžeme deformovat kruh, kouli i mnohoúhelníky a mnohostěny; pouze v nich nesmíme vytvořit dutinu (ani uzavřenou jako u míče, ani otevřenou jako u trubky). Při zkoumání různých souborů objektů, např. buněk ve tkáni, počtu zrn v krystalu či počtu stromů v lese, má μ_0/V význam hustoty nebo-li intenzity těchto objektů ve sledované oblasti, tj. počtu objektů na jednotku jejího obsahu.

Běžně známý je rozklad Eulerovy charakteristiky mnohostěnu M v $\mathbb{R}^3$: $\mu_0(M) = 1 = v - e + f - 1$, kde v, e, f jsou počty vrcholů, hran a stěn. Daleko méně známý je rozklad následující, platný i pro nekonvexní množiny:

$$\mu_0(A) = \sum_{i=0}^{\infty} (-1)^i \beta_i,$$

kde β_0 je počet izolovaných složek množiny A a β_i jsou počty nekontrahovatelných sfér $S^i(r)$ dimenze i či množin s nimi homeomorfních, které lze vepsat do A (poloměr r kontrahovatelné sféry můžeme neomezeně zmenšovat, aniž narušíme inkluzi $S^i(r) \subset A$); zřejmě $\beta_i = 0$ pro $i \geq d$ v $\mathbb{R}^d$. Jestliže $A \equiv S^i(r)$, je kontrahovatelná každá sféra dimenze $r < i$ (dva body $S^0(\rho)$ ležící na kružnici $S^1(r)$ či na libovolné sféře vyšší dimenze můžeme postupně stáhnout do jediného bodu stejně jako kružnici $S^1(\rho)$ ležící na sféře $S^2(r)$ či sféře libovolné vyšší dimenze atd.). Zřejmě je $\mu_0(A)$ homogenní řádu 0, tj. nezávisí na velikosti množiny A . Shodné hodnoty μ_0 bychom dostali také z podmínky aditivity (1). Hranicí úsečky L_1 jsou dva izolované body, odtud $\mu_0(\partial L_1) = 2$, kružnici C_1 složíme ze dvou polokružnic (jsou homeomorfní s úsečkou) majících společné dva body, odtud $\mu_0(C_1) = 0$, a povrch koule ∂K_3 složíme ze dvou polokoulí majících společnou kružnici, tj. $\mu_0(\partial K_3) = 2$.

Druhou důležitou aditivní invariantní valuací v $\mathbb{R}^d$ je délka, plocha, objem, obecně d -rozměrná Lebesgueova míra $\nu_d(K)$. Je homogenní řádu d a můžeme ji značit např. $\mu_d(K)$. Samozřejmě je rovna 0, jakmile dimenze K je menší než d ; valuaci s touto vlastností nazýváme *prostou*. V rámci systému valuací, který konstruujeme, ji budeme značit $\mu_d^d(K) \equiv \nu_d(K)$. Platí podobně základní věta jako pro μ_0 :

Věta 2 (Charakterizační teorém pro $\mu_d^d(K)$). *Nechť μ je spojitá pohybově invariantní prostá valuace nad konvexním okruhem $\mathcal{R}$. Pak existuje reálné c takové, že $\mu(A) = c\mu_d^d(A)$ pro všechna $A \in \mathcal{R}$.*

Pomocí těchto dvou elementárních valuací lze zavést řadu dalších valuací požadovaných vlastností. Konvexní množinu K dimenze d můžeme ortogonálně promítnout do lineárního podprostoru L_r^d , určit odpovídající Lebesgueovu míru této projekce $\nu_r(K|L_r^d)$, spočítat integrál těchto obsahů přes všechny možné orientace L_r^d a vydělit jej mírou těchto orientací, tj. $h_g(\mathcal{G}_r^d) = \gamma_r^d$. Interpretace tohoto funkcionálu je střední obsah projekce, tj. položíme

$$\mu_r^d(K) = \mathbf{E}\nu_r(K|L_r^d) = \frac{1}{\gamma_r^d} \int_{\mathcal{G}_r^d} \nu_r(K|L_r^d) d\sigma_r^d. \quad (2)$$

Lze dokázat, že takto zavedené funkcionály jsou aditivní a tedy valuace, navíc také spojité, monotónní, pohybově invariantní a homogenní řádu r . Pro $r = d$ nebo 0 je rovnice (2) triviální; $K|L_0^d$ je bod (počátek $\mathcal{O}$) pro každé $K \in \mathcal{K}$, takže $\nu_0(\mathcal{O}) = 1 = \mu_0(K)$ podle věty (1) (index d můžeme vynechat), a $K|L_d^d \equiv K$, takže rovnice (2) opět platí. Máme tak pro každou konvexní množinu v $\mathbb{R}^d$ celkem $(d+1)$ valuací výše uvedených vlastností. Platí pro ně následující stěžejní věta geometrické pravděpodobnosti (a také konvexní geometrie):

Věta 3 (Hadwigerův charakterizační teorém). *Systém valuací $\mu_0^d, \mu_1^d, \dots, \mu_d^d$ tvoří bázi vektorového prostoru všech spojitých pohybově invariantních valuací nad $\mathcal{K}$ v $\mathbb{R}^d$.*

To konkrétně znamená, že je-li $\phi(K)$ libovolná valuace těchto vlastností, pak musí platit

$$\phi(K) = \sum_{i=0}^d c_i \mu_i^d(K),$$

kde koeficienty c_i nezávisí na K . A navíc, když $\phi(K)$ je homogenní řádu r , pak z charakterizačního teorému plyne, že $\phi(K) = c_\phi \mu_r^d(K)$. Dále lze ještě dokázat, že je-li $\phi(K)$ monotónní, pak $c_i \geq 0$.

Vraťme se nyní k obsahu hranice množiny ∂K . Ten je homogenní pohybově invariantní spojitou valuací řádu $(d-1)$, a tedy musí platit $h_{d-1}(\partial K) =$

$c_{d-1}^d \mu_{d-1}^d(K)$. Koeficient c_{d-1}^d nezávisí na K a můžeme jej určit např. pro jednotkovou kouli dimenze d . Její obsah povrchu je $\mathcal{O}_{d-1}$ a její ortogonální projekce je jednotková koule $B_{d-1}(1)$ o obsahu κ_{d-1} . Odtud $c_{d-1}^d = \mathcal{O}_{d-1}/\kappa_{d-1}$. Dostali jsme velmi užitečný vztah Cauchyho

$$h_{d-1}(\partial K) = \frac{\mathcal{O}_{d-1}}{\kappa_{d-1}} \mathbf{E} \nu_{d-1}(K|L_{d-1}^d) = \frac{\mathcal{O}_{d-1}}{\kappa_{d-1}} \mu_{d-1}^d(K), \quad (3)$$

tj. obsah hranice libovolného konvexního tělesa je $\frac{\mathcal{O}_{d-1}}{\kappa_{d-1}}$ -násobkem obsahu jeho střední ortogonální projekce do nadroviny. Speciálně v $\mathbb{R}^3$ je to čtyřnásobek a v $\mathbb{R}^2$ je to π -násobek. To znamená, že známý vztah mezi obsahem rovinné projekce koule a obsahem jejího povrchu resp. mezi obvodem a průměrem kruhu platí **zcela obecně** pro všechny konvexní množiny odpovídající dimenze. Takže např. obsah střední rovinné projekce hranolu o hranách a, b, c je $(ab + ac + bc)/2$ a válce o poloměru ρ a výšce h je $\frac{\pi}{2}r(r + h)$.

Důležitá je rovněž valuace $\mu_1^d(K)$, tj. střední obsah ortogonální projekce K do všech L_1^d . Nazývá se *střední šířka* nebo také *střední Feretův* či *klešťový průměr* – oba názvy jsou vzaty z technické praxe a odpovídají střednímu údaji mikrometru při různých otočeních měřeného předmětu.

Jako čtyři charakteristiky konvexního tělesa tvořící bázi valuací v $\mathbb{R}^3$ obvykle volíme jeho objem ($V(K) = \mu_3^3(K) = \nu_3(K)$), obsah hranice $S(\partial K) \equiv S(K)$, střední šířku (značíme $w(K) = \mu_1^3(K)$) a Eulerovu charakteristiku. Podobně v $\mathbb{R}^2$ volíme obsah (tedy opět Lebesgueovu míru) množiny $A(K)$, častěji délku jejího obvodu $L(\partial K)$ než střední šířku w , a Eulerovu charakteristiku. V $\mathbb{R}^1$ bázi tvoří délka intervalu a Eulerova charakteristika.

V $\mathbb{R}^2$ je $\mu_{d-1}^d(K) \equiv \mu_1^2(K)$ a tedy podle rovnice (3) je obvod konvexního obrazce $L(\partial K) = \pi w(K)$. Relace platí opět nejen pro kruh, ale zcela obecně; $w = 4a/\pi$ platí nejen pro čtverec a kosočtverec o stranách a , ale i pro limitní případ smykově deformovaných obrazců na úsečku délky $L(\partial K)/2$. Odtud je střední šířka úsečky délky L rovna $\frac{2L}{\pi}$, jak odvodíme přímo níže. Ze stejného důvodu je střední obsah projekce plošného obrazce plochy A v $\mathbb{R}^3$ roven $A/2$, protože jej můžeme považovat výsledek limitního smyku tělesa s obsahem hranice $2A$, pro nějž také musí platit rovnice (3).

Pro snadnou práci v obecně d -rozměrném prostoru bývají valuace různě normovány. Velmi často se používají tzv. *Minkowského funkcionály* či *integrály příčného průřezu* (v angličtině i němčině shodně *quermassintegrale*) $W_r^d(K)$. Jsou spojitě, pohybově invariantní a homogenní řádu $(d - r)$ a normovány tak, že pro jednotkovou kouli jsou všechny rovny jejímu objemu, tj. κ_d . Z charakterizačního teoremu pak plyne, že musejí být úměrné μ_{d-r}^d , jež

je pro jednotkovou kouli rovno κ_{d-r} , takže

$$W_r^d(K) = \frac{\kappa_d}{\kappa_{d-r}} \mu_{d-r}^d(K). \quad (4)$$

Rovnice (4) se nazývá zobecněný vztah Cauchyho; jeho speciální případ jsme dostali výše. $W_0^d(K)$ udává Lebesgueovu míru množiny (její objem), $dW_1^d(K)$ je obsah její hranice a $\mu_1^d(K) = (\kappa_{d-1}/\kappa_d)W_{d-1}^d(K)$ je její střední šířka.

Jestliže je dimenze K obecně $s \leq d$ (pak budeme psát explicitně K_s), potom prvních $(d-s)$ Minkowského funkcionálů $W_i^d(K_s)$, $i = 0, \dots, d-s-1$, právě tak jako posledních $(d-s)$ valuací $\mu_j^d(K_s)$, $j = s, \dots, d$ je rovno nule. Nenulové funkcionály $W_i^d(K_s)$, $d \geq i \geq d-s$, lze vyjádřit pomocí $W_j^s(K_s)$ prostoru $\mathbb{R}^s$ podle vztahu

$$W_i^d(K_s) = \frac{\mathcal{O}_{i+s-d}}{\mathcal{O}_i} \frac{\mathcal{O}_{d-s}\kappa_{d-s}}{2\binom{d}{s}} W_{i+s-d}^s(K_s),$$

takže pro úsečku délky L , tj. $s = 1$, je $W_{d-1}^d(K_1) = \kappa_{d-1}L/d$ a $W_d^d(K_1) = \kappa_d$. Odtud s použitím vztahu (4) platí

$$\mu_r^d(K_s^d) = \frac{\kappa_r}{\kappa_d} W_{d-r}^d(K_s) = \frac{\mathcal{O}_{s-r}\mathcal{O}_d\kappa_r}{\mathcal{O}_{d-r}\mathcal{O}_s\kappa_s} W_{s-r}^s(K_s). \quad (5)$$

Speciálně je např. střední délka ortogonální projekce úsečky K_1 délky L v $\mathbb{R}^d$

$$\mathbf{E}h_1(K_1|L_1^d) = \mu_1^d(K_1) = \frac{\mathcal{O}_d}{\pi\mathcal{O}_{d-1}}L,$$

tj. $2L/\pi$ v $\mathbb{R}^2$ (tento výsledek jsme ostatně dostali již v Buffonově úloze o jehle) a $1/2$ v $\mathbb{R}^3$.

6. Croftonův vztah

Až dosud jsme pracovali s projekcemi konvexních množin. Přejdeme nyní k jejich průnikům s r -rovinami F_r^d . Ty jsou určeny svou orientací, tedy bodem $\mathcal{G}_r^d$, a dále svým bodovým průnikem s komplementárním ortogonálním lineárním podprostorem $L_{r\perp}^d \equiv L_{d-r}^d$. Každé r -rovině tedy odpovídá jeden bod na určité varietě dimenze $(d-r)(r+1)$, kterou označíme $\mathcal{F}_{d,r}$ a elementární míru (nazývá se *kinematická míra*) na její infinitesimální oblasti označíme $d\phi(F_r^d \uparrow K)$. Varieta, na rozdíl od $\mathcal{G}_r^d$ a $\mathcal{S}_r^d$ zřejmě není omezená, taková je pouze ta její oblast, která popisuje neprázdné průniky $F_r^d \cap K$ s omezenou

množinou K . Obsah této oblasti budeme značit $\phi_r^d(K)$. Vytvořme následující integrál:

$$\frac{1}{\gamma_r^d} \int_{\mathcal{F}_{d,r}} W_i^r(K \cap F_r^d) d\phi(F_r^d \uparrow K), \quad r \geq i.$$

Zřejmě $K \cap F_r^d$ je pro $F_r^d \uparrow K$ oblast v F_r^d a tedy $W_i^r(K \cap F_r^d)$ má smysl pro $0 \leq i \leq r$. O integrálu lze dokázat, že je pohybově invariantní, monotónní a homogenní valuace. Řád homogenity $W_i^r(K \cap F_r^d)$ je $(r - i)$, zatím co řád integračního oboru odpovídá projekci $K|L_{d-r}^d$ a je tedy $(d - r)$. Celkem je tedy integrál homogenní řádu $(d - i)$ a podle charakterizačního teorému musí být úměrný $W_i^d(K)$. Dostaneme (jeden z mnoha) *Croftonových vztahů*

$$W_i^d(K) = \frac{c_{ir}^d}{\gamma_r^d} \int_{\mathcal{F}_{d,r}} W_i^r(K \cap F_r^d) d\phi(F_r^d \uparrow K), \quad r \geq i, \quad (6)$$

kde c_{ir}^d je konstanta nezávislá na K . Dosazením jednotkové koule za K lze odvodit $c_{ir}^d = \frac{\kappa_d \kappa_{r-i}}{\kappa_r \kappa_{d-i}}$.

Všimněme si Croftonova vztahu pro $i = r$. Zřejmě $W_r^r(K \cap F_r^d) = \kappa_r$ pro průniky neprázdné, a nule jinak a $c_{rr}^d = \frac{\kappa_d}{\kappa_r \kappa_{d-r}}$. S použitím rovnice (4) dostaneme tak $\gamma_r^d \mu_{d-r}^d = \phi_r^d(F_r^d \uparrow K) = \phi_r^d(K)$, takže výraz na levé straně této rovnice je právě míra souboru všech r -rovin zasahujících K . Konkrétně v $\mathbb{R}^3$ je $\phi_2^3(K) = 2\pi w(K)$ (i pro K dimenze nižší než d) a $\phi_1^3(K_3) = \frac{\pi}{2} S(K_3)$, v $\mathbb{R}^2$ je $\phi_1^2(K) = \pi w(K) = L(\partial K)$.

Odtud také plyne pro množinu $K' \in K$, že

$$P(F_r^d \uparrow K' | F_r^d \uparrow K) = \frac{\phi_r^d(K')}{\phi_r^d(K)} = \frac{\mu_{d-r}^d(K')}{\mu_{d-r}^d(K)},$$

což je obecný tvar Sylvestrova vztahu, jehož speciální tvar pro $r = 0, 2$ jsme dostali v kap. 3.

Rovnici (6) můžeme tedy zapsat ve tvaru

$$W_i^d(K_d) = \frac{\kappa_d \kappa_{r-i}}{\kappa_r \kappa_{d-i}} \mu_{d-r}^d(K_d) \mathbf{E} W_i^r(K_d \cap F_r^d), \quad r \geq i. \quad (7)$$

Důležitost předchozích poznatků demonstrují následující dva příklady.

Příklad 11. Určete míru $\phi_2^3(K)$ rovin zasahujících uzavřenou krychlovou oblast K v $\mathbb{R}^3$ objemu $V = a^3$.

Obecně je $\phi_1^3(K) = 2\pi w(K)$ (pro roviny s normálou $\mathbf{n}$ je to právě délka projekce $K|\mathbf{n}$, tj. šířka $w(K, \mathbf{n})$; při uvažování všech orientací zasahujících K je $w(K, \mathbf{n})$ nahrazeno svou střední hodnotou $w(K)$ a vynásobeno mírou

všech orientací normály $\mathbf{n}$, tj. 2π). 2π -násobek střední šířky $w(K)$ je obecně roven integrálu střední křivosti tělesa, jehož hodnota pro obecný mnohostěn je rovna $\frac{1}{2} \sum_i (\pi - \alpha_i) a_i$, kde α_i jsou vnitřní úhly stěn a a_i jsou délky odpovídajících hran. Odtud pro krychli K je $2\pi w = 6\frac{\pi}{2}a = 3\pi a$.

Příklad 12. V krychli z předcházejícího příkladu jsou rovnoměrně náhodně rozmístěny dva stejně početné soubory A , B kulových částic a_i , b_i , $i = 1, \dots, N$, o průměrech $d_A = 2d_B$. Jaká je očekávaná hodnota poměru počtu zasažených částic ze souborů A , B náhodně vybranou rovinou $F \uparrow K$?

Míry rovin zasahujících krychli, částici a_i a b_i jsou s ohledem na předcházející příklad $3\pi a$, $2\pi d_A$, πd_A . Odtud je podmíněná pravděpodobnost $P(F \uparrow a_i | F \uparrow K)$ rovna $\frac{2}{3}d_A/a$ a je poloviční pro částice b_i . V rovině F lze tedy očekávat dvojnásobný počet zasažených částic typu A než B .

Rovina si tedy vybírá částice s váhou úměrnou jejich střední šířce. Podobně bychom dostali, že náhodná přímka vybírá částice s váhou úměrnou střednímu obsahu jejich rovinné projekce, náhodný bod úměrně jejich objemu. Náhodný výběr geometrickým prostředkem „průnik“ je tedy - na rozdíl od geometrického výběrového prostředku „projekce“ - **vážený!** Intuitivně jsme těchto pravidel používali v historických úlohách, zde je jejich obecná formulace.

Druhou důležitou vlastností geometrického výběru prováděného r -rovinami je ztráta informace nepřímo úměrná dimenzi sondy, která je specifikována podmínkou $r \geq i$. Konkrétně d -rozměrný obsah $V(K) \equiv W_0^d(K)$ lze odhadnout z řezů $F_r^d \cap K$ při $d \geq r \geq 0$, obsah hranice již jen z řezů pro $d \geq r \geq 1$ a k odhadu Eulerovy charakteristiky souboru částic je obecně nezbytná d -rozměrná oblast.

Poslední obecný vztah, který zde uvedeme, je *Steinerova věta* pro vnitřní a vnější rovnoběžnou množinu, s níž jsme se setkali v kapitole 2.

Věta 4. *Nechť K je konvexní množina. Pak Minkowského funkcionály vnější rovnoběžné množiny K_ρ jsou*

$$W_i^d(K_\rho) = \sum_{k=0}^{d-i} \binom{d-i}{k} W_{k+i}^d(K) \rho^k, \quad (8)$$

Jestliže existuje konvexní množina A taková, že $K = A_{\rho'}$, pak pro $\rho \leq \rho'$ jsou Minkowského funkcionály vnitřní rovnoběžné množiny $K_{-\rho}$

$$W_i^d(K_{-\rho}) = \sum_{k=0}^{d-i} \binom{d-i}{k} W_{k+i}^d(K) (-\rho)^k. \quad (9)$$

Speciálně v $\mathbb{R}^3$ pro konvexní množinu K objemu V s obsahem povrchu S a střední šířkou w je objem $V(K_\rho) = V + \rho S + 2\pi\rho^2 w + \frac{4}{3}\pi\rho^3$. Vztah platí i pro konvexní množiny dimenze nižší než 3; např. vnější rovnoběžné těleso H_ρ úsečky H délky h je sjednocení válce se dvěma polokoulemi (v místě podstav) o objemu $\pi\rho^2 h + \frac{4}{3}\pi\rho^3$, což dostaneme, dosadíme-li správně za střední šířku úsečky v $\mathbb{R}^3$ hodnotu $\frac{1}{2}h$ odvozenou výše. Podobně v $\mathbb{R}^2$ pro konvexní obrazec plochy A a obvodu P je $A(K_\rho) = A + \rho P + \pi\rho^2$.

Dosud jsme většinou uvažovali jedinou konvexní množinu. Valuační jsme však zavedli proto, abychom je mohli používat i pro konečná sjednocení množin. Proto prakticky všechny uvedené výsledky jsou aplikovatelné i v rámci konvexního okruhu $\mathcal{R}$. Poněkud se mění pouze interpretace jednotlivých valuací, vedle projekcí v běžném smyslu slova se zavádějí tzv. totální projekce, jejichž obsah je součtem obsahů jednotlivých překrývajících se množin atd. Např. obsah projekce činky ve směru její osy je dvojnásobek obsahů projekcí jejich koulí. Důležitou skutečností je především to, že Hadwigerův charakterizační teorém platí i pro množiny konvexního okruhu a odtud plyne i obecnější platnost rovnic (6) a (7).

K čemu jsou uvedené vztahy užitečné? Jak jsme viděli v příkladu 12, uvedené vztahy pro průniky r -rovinami (např. rovnice (6, 7)) lze aplikovat i na množiny $Y \in K$. Každá $F_r^d \uparrow Y$ je totiž $F_r^d \uparrow K$, a jestliže naopak $F_r^d \uparrow K$ míjí Y , jsou Minkovského funkcionály $W(F_r^d \cap Y)$ rovny nule a k hodnotám integrálů nepřispívají. Takže můžeme psát také

$$W_i^d(Y) = c_{ir}^d \mu_{d-r}^d(K) \mathbf{E} W_i^r(Y \cap F_r^d), \quad (10)$$

kde střední hodnota se teď vztahuje ke všem $F_r^d \uparrow K$! Mějme reálný vzorek K obsahující částice či dutiny či obojí Y_d (materiál se zpevňujícími částicemi či vlákny, ementálský sýr nebo vánočka s rozinkami). V takových případech se především zajímáme o poměr objemu složky Y k celkovému objemu kompozitního objektu - tzv. objemový podíl $V_V(Y) = V(Y)/V(K)$. Zapsáním rovnice (10) pro Y a K v $\mathbb{R}^3$ a jejich vydělením dostaneme

$$V_V(Y) = \frac{\mathbf{E} W_0^r(Y_d \cap F_r^3)}{\mathbf{E} W_0^r(K_3 \cap F_r^3)}, \quad (11)$$

kde W_0^r má význam plošného obsahu při $r = 2$, délky při $r = 1$ a počtu zasahujících bodů při $r = 0$. Pokud udržujeme jmenovatel konstantní (měření provádíme ve stále stejném „pozorovacím okénku“ plošného obsahu A či na úsečkách stejné celkové délky L plně ležících v K nebo konečně na pevně zvoleném souboru $\mathcal{N}$ náhodně či pravidelně rozložených N bodů, které vždy

všechny zasáhnou K_3), pak můžeme psát

$$V_V(Y) = \mathbf{E}A(Y \cap F_2^3)/A = \mathbf{E}L(Y \cap F_1^3)/L = \mathbf{E}\kappa_0(Y \cap F_0^3)/N. \quad (12)$$

Nejsou-li tyto podmínky splněny, musíme respektovat, že se jedná o poměr středních hodnot a nikoliv o střední hodnotu poměru. Podobně můžeme odhadovat také např. poměr $S_V(Y) = S(Y)/V(k)$, udávající poměr plošného obsahu hranic souboru Y k objemu celého vzorku K . Výše zmíněná ztráta informace má za následek, že počet izolovaných komponent v souboru Y lze obecně odhadovat pouze z d -rozměrné podmnožiny (vzorku) Y přímým spočítáním komponent a korekcí na okrajové jevy. Podobor geometrické pravděpodobnosti a statistiky nazývaný *stereologie* řeší problematiku odhadu geometrických charakteristik nejrozličnějších jednoduchých i složitých objektů z geometrických výběrů s použitím pravděpodobnostních relací typu Croftonova či Sylvestrova vztahu resp. Cauchyho relace aj.

7. Závěrečné poznámky

Probrán byl jen velmi úzký okruh úloh geometrické pravděpodobnosti, a to jen v jejich nejjednodušší formě. Buffonovu úlohu o jehle lze zobecnit nejen na případ dvou systémů křivek, jak bylo naznačeno v kap. 2, ale také na problém interakce křivek a ploch v $\mathbb{R}^3$. Problematika interakcí v rámci konvexního okruhu byla jen naznačena, v současné době již však probíhá její rozšíření do podstatně obecnějších množinových systémů.

Za zmínku stojí alespoň jedna z nejslavnějších úloh – *Sylvestrův problém*. V něm je hledána pravděpodobnost, že čtyři nezávisle náhodně umístěné body v obrazci Z vytvoří konvexní čtyřúhelník. Řešení je závislé na tvaru oblasti, ale nemění se při její afinní transformaci. Pravděpodobnost je maximální pro elipsy (0.71) a minimální pro trojúhelníky (0.67).

Obsažná je oblast geometrických nerovností mezi valuacemi, jejich soubor vydá na celou knihu [1]. Nejznámější z nich jsou klasická rovinná izoperimetrická nerovnost $P^2 - 4\pi A \geq 0$ mezi obsahem A a obvodem P konvexního obrazce a její prostorová obdoba $S^3 - 36\pi V^2 \geq 0$, podle nichž kruh a koule mají největší obsah při daném obsahu hranice (analogické nerovnosti platí ve všech dimenzích: $S^d - d^d \kappa_d V^{d-1}$). Méně již je známo, že podobná (Bierbachova) nerovnost platí i pro střední šířky:

$$\kappa_d w^d - 2^d V \geq 0.$$

Atraktivní je problematika vzájemných vzdáleností mezi složkami geometrických soustav. Jejich dva základní typy jsou lineární orientované a sférické „izotropní“ vzdálenosti. U prvního typu je nejznámější lineární vzdálenost

mezi izolovanými „částicemi“ c_i souboru $C = \bigcup_i c_i$ měřená na náhodně vybraných testovacích přímkách F_1^d : na nich se střídají tětivy částic $F_1^d \cap c_i$ s „tětivami“ komplementu C^{com} , jež v jistém smyslu vypovídají o vzdálenostech částic. Bohužel však jejich střední hodnota je na rozložení částic nezávislá a nezávisí dokonce ani na tvaru testovací čáry (jedná se vlastně o prostorovou analogii Buffonovy úlohy o jehle, v níž testovací čára odhaduje obsah sjednocení povrchů částic). Přesto je tento typ „vzdálenosti“ v praxi využíván jako charakteristika rozmístění částic. Sférická vzdálenost se obvykle určuje od náhodného bodu x komplementu C^{com} a je rovna poloměru sféry se středem v x , která se právě dotýká hranice ∂c_i částice nejbližší bodu x . Nazývá se *vnější (externí) sférická kontaktní vzdálenost* a lze ji použít i k charakterizaci bodových soustav (procesů). Pokud za bod x zvolíme také bod procesu, dostáváme *vzdálenost nejbližších sousedů*. Teorie vzdáleností v různých soustavách objektů je velmi podrobně rozpracována a široce využívána v praxi.

Poznatky, které jsme probrali pro eukleidovský prostor, mohou být zobecněny i na prostory zakřivené, speciálně na d -rozměrné sféry $\mathbb{S}^d$. Pro potřeby automatické obrazové analýzy je zase nezbytné řešit podobnou problematiku v diskretních bodových prostorech (pravidelných mřížkách obrazových bodů).

Žádné dva lipové listy nejsou stejné, ale všechny se navzájem podobají. Podle tvaru bezpečně poznáme jablko od hrušky, určíme plemena psů, poznáme rostliny podle řady tvarových znaků atd. Ve všech případech se jedná o realizace omezeného náhodného objektu „břečtan“, „leopard“, „křemen“. V lese a na záhonu jsou rostlinné realizace rozmístěny v systematicky či náhodně, v minerálu jsou vedle sebe bez překrývání uspořádány prostor beze zbytku vyplňující krystaly, a obec i státy a světadíly jsou „rozparcelovány“ podobným způsobem. Můžeme tedy mluvit o realizacích více či méně náhodných procesů, jejichž komponenty mají definované náhodné rozměry a tvary a zákonitě náhodně jsou rozmístěny v prostoru. Touto problematikou se zabývá část teorie geometrické pravděpodobnosti nazývaná *stochastickou geometrií*. V jejím rámci byla vyvinuta řada nástrojů - stochastických funkcionalů - pro popis náhodných procesů bodů, vláken, ploch a částic. Systematicky je rozvíjena i velmi obtížná problematika náhodných tvarů a jejich popisu. Všimněme si alespoň tří velmi jednoduchých modelů takových procesů.

V prvním z nich jsou jednotlivými objekty body, odpadá tedy náhodnost rozměru a tvaru. Jestliže jsou body v $\mathbb{R}^d$ rozmístěny nezávisle rovnoměrně náhodně, nazývá se model *Poissonův bodový proces* (PBP). Můžeme jej charakterizovat několika způsoby. Např. tím, že počty bodů n v libovolné omezené borelovské množině obsahu V mají Poissonovo rozdělení $P(n) =$

$\frac{(\lambda V)^n}{n!} \exp(-\lambda V)$, kde λ je intenzita bodového procesu. Stačí však zadat pouze tzv. pravděpodobnost dutiny, tj. pravděpodobnost $P_K(n=0) = \exp(-\lambda V)$, že v kompaktní množině K obsahu V není žádný bod procesu. Poissonův bodový proces hraje v teorii procesů podobně základní roli jako normální rozdělení u náhodných veličin; pravidelnější rozložení bodů mají tzv. uspořádané procesy s negativní interakcí (body se do jisté míry odpuzují), méně pravidelné jsou naopak procesy shlukové (body se přitahují).

V modelu *zárodek-zrno* je zadán libovolný bodový proces zárodků a do nich je (obvykle nezávisle) umístěna náhodná množina - zrno (např. koule se zadaným rozdělením poloměrů). Jestliže proces zárodků je PBP, nazývá se model *Booleovým*. Jeho problémem je však to, že v PBP dochází v jisté míře ke shlukování a již při nízkém objemovém podílu zrn (nad 5%) dochází k překrývání zrn, takže jeho použití je bezesporne pouze pro dutiny (model chleba, spěkaných kompozitních materiálů, travertinu atd.)

Poslední model popisuje systém množin, který bez překrývání vyplňuje celý prostor nebo alespoň jeho omezenou oblast. Výchozím je opět libovolný bodový proces. Každému jeho bodu - generátoru - je přiřazena oblast zvaná cela, sjednocující všechny body obklopujícího prostoru, jejichž vzdálenost od jiných generátorů je menší resp. nejvýše stejná. Hranice cel jsou tedy tvořeny body, jež jsou stejně vzdálené od více generátorů. Model se nazývá *Voronoi-ova* (jestliže proces generátorů je PBP, pak *Poissonova-Voronoi-ova*) *teselace*, v rovině často také *mozaika*.

Doporučená literatura uvádí u nás dostupné základní monografie, jejichž názvy vesměs zřetelně vymezují jejich obsah.

Literatura

- [1] Burago Yu.D. a Zalgaller V.A. (1980). *Geometric Inequalities*. Springer-Verlag, Berlin.
- [2] Kendall M. a Moran P. (1963). *Geometric probability*. Hafner Publishing Company, New York.
- [3] Klain D.A. a Rota G.-C. (1997). *Introduction to Geometric Probability*. Cambridge University Press.
- [4] Santaló L.A. (1976). *Integral Geometry and Geometric Probability*. Addison-Wesley, Reading (Mass.).
- [5] Saxl I. (1989). *Stereology of Objects with Internal Structure*. Academia, Prague & Elsevier, Amsterdam.

- [6] Saxl I., Pelikán K., Rataj J. a Besterčí M. (1995). *Quantification and Modelling of Heterogeneous Systems*. Cambridge Int. Science Publishing, Cambridge.
- [7] Serra J. (1982). *Image Analysis and Mathematical Morphology*. Academic Press, London.
- [8] Solomon H. (1978). *Geometric Probability*. Society for Industrial and Applied Mathematics. Philadelphia.
- [9] Stoyan D., Kendall W.S. a Mecke J. (1995). *Stochastic Geometry and its Applications*. J. Wiley & Sons, New York.
- [10] Stoyan D. a Stoyan H. (1994). *Fractals, Random Shapes and Point Fields*. J. Wiley & Sons, Chichester.
- [11] Weibel E. (1995). *Stereological Methods, Vol. 2*. Academic Press, London.

NÁHODNÉ PROCESY

Viktor Beneš

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

viktor.benes@mff.cuni.cz

1. Úvod

Náhodné procesy představují modely dynamických stochastických dějů, které lze popsat systémem náhodných veličin měnících se v čase nebo v prostoru. Tyto náhodné veličiny zpravidla nejsou nezávislé a existují různé třídy modelů popisující jejich vzájemnou vazbu. Vedle stacionárních procesů (Anděl, 1976), kde požadavkem je, aby se pravděpodobnostní charakteristiky neměnily v čase či prostoru, existuje třída Markovských procesů s krátkým dosahem závislosti. Předložená přednáška se omezuje právě na tento model oblíbený pro svou jednoduchost a přitom širokou aplikaci.

2. Markovovy řetězce s konečnou množinou stavů

Definice 1. Nechtě $S = \{s_1, \dots, s_k\}$ je konečná množina, $\mathbb{P}$ je matice typu $k \times k$ s nezápornými prvky p_{ij} splňující $\sum_{j=1}^k p_{ij} = 1$ pro každé $i = 1, \dots, k$. Posloupnost náhodných veličin $X = (X_0, X_1, \dots)$ se nazývá homogenní Markovův řetězec s přechodovou maticí $\mathbb{P}$ jestliže pro každé n přirozené a pro každá $i, j, i_{n-1}, \dots, i_0$ z $\{1, \dots, k\}$ platí

$$\begin{aligned} P(X_{n+1} = s_j \mid X_n = s_i, X_{n-1} = s_{i_{n-1}}, \dots, X_0 = s_{i_0}) = \\ = P(X_{n+1} = s_j \mid X_n = s_i) = p_{ij}, \end{aligned} \quad (1)$$

kde P značí pravděpodobnost.

Přívlastek homogenní, ze kterého vyplývá, že p_{ij} v (1) nezávisí na n , budeme dále vynechávat. Rovnosti (1) se říká Markovská vlastnost. Lze ji interpretovat tak, že k předpovědi toho, co se stane v čase $n+1$, stačí uvažovat

Klíčová slova: Markovský řetězec, stacionární rozdělení.

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839.

co se stane v čase n , zatímco časové okamžiky $0, \dots, n-1$ nedávají žádnou doplňující informaci. Toto je vlastnost modelů, kterými se budeme zabývat v této kapitole.

Vedle přechodové matice je další charakteristikou Markovova řetězce počáteční rozdělení $\xi^{(0)}$, píšeme-li

$$\xi^{(n)} = (\xi_1^{(n)}, \dots, \xi_k^{(n)}),$$

kde $\xi_i^{(n)} = P(X_n = s_i)$.

Věta 1. Platí $\xi^{(n)} = \xi^{(0)}\mathbb{P}^n$ pro každé n přirozené.

Důkaz: Matematickou indukcí. Pro $n = 1$ vyjdeme ze vztahu

$$P(X_1 = s_j) = \sum_{i=1}^k P(X_1 = s_j \mid X_0 = s_i)P(X_0 = s_i),$$

který představuje právě j -tý prvek vektoru $\xi^{(1)} = \xi^{(0)}\mathbb{P}$. Indukční krok užívá stejný argument: je-li $\xi^{(n)} = \xi^{(0)}\mathbb{P}^n$, potom $\xi^{(n+1)} = \xi^{(n)}\mathbb{P} = \xi^{(0)}\mathbb{P}^n\mathbb{P} = \xi^{(0)}\mathbb{P}^{n+1}$.

Prvky mocniny $\mathbb{P}^n$ značíme $p_{ij}^{(n)}$, vzhledem k tvrzení věty 1 je lze interpretovat jako pravděpodobnosti přechodu řetězce ze stavu i do stavu j v n krocích.

Definice 2. Řekneme, že stav s_j je dosažitelný ze stavu s_i , jestliže existuje n přirozené tak, že $P(X_n = s_j \mid X_0 = s_i) > 0$. Markovův řetězec je nerozložitelný, je-li libovolný stav dosažitelný z libovolného jiného stavu. Perioda $d(s_i)$ stavu s_i je definována jako největší společný dělitel množiny $\{n \geq 1; p_{ii}^{(n)} > 0\}$. Je-li $d(s_i) = 1$ nazývá se stav s_i aperiodický. Markovův řetězec je aperiodický jsou-li všechny jeho stavy aperiodické. Jinak se nazývá periodický.

Příklad 1. Zjednodušený model počasí se dvěma stavy $s_1 = \text{„deštivo“}$, $s_2 = \text{„slunečno“}$ je založen na předpokladu, že předpověď na zítřejší den se odvíjí od dnešního dne. Uvažme např. přechodovou matici

$$\mathbb{P} = \begin{pmatrix} 0,6 & 0,4 \\ 0,3 & 0,7 \end{pmatrix}.$$

řetězec s touto přechodovou maticí je zřejmě nerozložitelný a aperiodický neboť všechny prvky matice $\mathbb{P}$ jsou nenulové.

Příklad 2. Nezávisle házíte hrací kostkou, výsledky Y_i , $i = 0, 1, \dots$ představují počet bodů. Tedy stavový prostor $S = \{1, 2, 3, 4, 5, 6\}$. Položme $X_0 = Y_0$, dále $X_{i+1} = \max\{X_i, Y_{i+1}\}$, $i = 0, 1, \dots$ Markovův řetězec $X = (X_0, X_1, \dots)$ na S není nerozložitelný, neboť ze stavu 6 není dosažitelný žádný jiný stav než opět 6. říkáme, že X je rozložitelný.

Příklad 3. Uvažme náhodnou procházku ve vesnici se čtyřmi ulicemi, které tvoří strany čtverce. Chodec ve vrcholu čtverce jde s pravděpodobností $1/2$ vlevo nebo vpravo. Označme vrcholy $S = \{1, 2, 3, 4\}$, matice pravděpodobností přechodu má tvar

$$\mathbb{P} = \begin{pmatrix} 0 & \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{2} & 0 & \frac{1}{2} & 0 \\ 0 & \frac{1}{2} & 0 & \frac{1}{2} \\ \frac{1}{2} & 0 & \frac{1}{2} & 0 \end{pmatrix}. \quad (2)$$

Zřejmě všechny stavy jsou periodické s periodou 2 a Markovův řetězec náhodné procházky je periodický.

Věta 2. *Nechť X je nerozložitelný aperiodický Markovův řetězec, potom existuje N přirozené tak, že $p_{ij}^{(n)} > 0$ pro každé $i, j \in \{1, \dots, k\}$ a každé $n \geq N$.*

Důkaz: Nejdříve ukážeme tvrzení v případě $i = j$. Pro $s_i \in S$ buď $A_i = \{n \geq 1; p_{ii}^{(n)} > 0\}$. Z aperiodicity plyne, že největší společný dělitel A_i je roven jedné. Dále ukážeme, že A_i je uzavřená vůči sčítání. Tyto dvě vlastnosti již podle známého lemmatu z teorie čísel (Brémaud, 1998, Appendix) mají za následek existenci M_i přirozeného tak, že $p_{ii}^{(n)} > 0$ pro každé $n \geq M_i$. Buďte tedy $n, m \in A_i$, je $P(X_n = s_i | X_0 = s_i) > 0$ a $P(X_{n+m} = s_i | X_n = s_i) > 0$. Odsud s využitím Markovské vlastnosti

$$\begin{aligned} P(X_{n+m} = s_i | X_0 = s_i) &\geq P(X_{n+m} = s_i, X_n = s_i | X_0 = s_i) = \\ &= P(X_{n+m} = s_i | X_n = s_i)P(X_n = s_i | X_0 = s_i) > 0, \end{aligned}$$

takže $n + m \in A_i$, což je uzavřenost A_i vůči sčítání. Položme $M = \max\{M_1, \dots, M_k\}$.

Buď nyní $i \neq j$, z nerozložitelnosti plyne existence n_{ij} přirozeného tak, že $p_{ij}^{(n_{ij})} > 0$. Položme $N_{ij} = M + n_{ij}$. Pro každé $m > N_{ij}$ je podobně jako výše

$$\begin{aligned} P(X_m = s_j | X_0 = s_i) &\geq P(X_m = s_j, X_{m-n_{ij}} = s_i | X_0 = s_i) = \\ &= P(X_m = s_j | X_{m-n_{ij}} = s_i)P(X_{m-n_{ij}} = s_i | X_0 = s_i) > 0. \end{aligned}$$

Tvrzení věty nyní plyne pro $N = \max\{N_{ij}; i, j \in \{1, \dots, k\}\}$.

3. Klasifikace stavů Markovova řetězce

Definice 3. Nechť Markovův řetězec X s přechodovou maticí $\mathbb{P}$ startuje ve stavu s_i , potom definujeme čas dosažení stavu s_j předpisem

$$\tau_{ij} = \min\{n \geq 1; X_n = s_j\},$$

přičemž $\tau_{ij} = \infty$ jestliže řetězec nikdy nedosáhne s_j .

Speciálně pro $i = j$ se τ_{ii} nazývá čas návratu.

Věta 3. Pro nerozložitelný aperiodický Markovův řetězec se stavovým prostorem $S = \{s_1, \dots, s_k\}$ a přechodovou maticí $\mathbb{P}$ platí pro libovolné i, j

$$P(\tau_{ij} < \infty) = 1. \quad (3)$$

Navíc střední čas dosažení je konečný, tj.

$$E\tau_{ij} < \infty.$$

Důkaz: Z věty 3 můžeme najít $m < \infty$ tak, že $p_{ij}^{(m)} > 0$ pro každé $i, j \in \{1, \dots, k\}$. Zvolme takové m , položme $\alpha = \min\{p_{ij}^{(m)}; i, j \in \{1, \dots, k\}\}$, je tedy $\alpha > 0$. Dále zvolme pevná i, j , nechť řetězec startuje ze stavu s_i . Platí $P(\tau_{ij} > m) \leq P(X_m \neq s_j) \leq 1 - \alpha$. Dále

$$\begin{aligned} P(\tau_{ij} > 2m) &= P(\tau_{ij} > m)P(\tau_{ij} > 2m \mid \tau_{ij} > m) \\ &\leq P(\tau_{ij} > m)P(X_{2m} \neq s_j \mid \tau_{ij} > m) \\ &\leq (1 - \alpha)^2. \end{aligned}$$

Iterujeme tento argument a dostáváme pro každé l

$$\begin{aligned} P(\tau_{ij} > lm) &= P(\tau_{ij} > m)P(\tau_{ij} > 2m \mid \tau_{ij} > m) \dots \\ &\quad \times P(\tau_{ij} > lm \mid \tau_{ij} > (l-1)m) \\ &\leq (1 - \alpha)^l, \end{aligned}$$

což konverguje k 0 když $l \rightarrow \infty$. Tedy $P(\tau_{ij} = \infty) = 0$ a (3) je dokázáno. Dále

$$\begin{aligned} E\tau_{ij} &= \sum_{n=0}^{\infty} P(\tau_{ij} > n) = \sum_{l=0}^{\infty} \sum_{n=lm}^{(l+1)m-1} P(\tau_{ij} > n) \\ &\leq \sum_{l=0}^{\infty} \sum_{n=lm}^{(l+1)m-1} P(\tau_{ij} > lm) = m \sum_{l=0}^{\infty} P(\tau_{ij} > lm) \\ &\leq m \sum_{l=0}^{\infty} (1 - \alpha)^l = \frac{m}{\alpha} < \infty. \end{aligned}$$

Definice 4. Stav s_i se nazývá trvalý, jestliže $Pr(\tau_{ii} < \infty) = 1$. V opačném případě, tj. jestliže $Pr(\tau_{ii} = \infty) > 0$, stav s_i se nazývá přechodný. Trvalý stav s_j se nazývá nenulový, jestliže $E\tau_{jj} < \infty$, v opačném případě se nazývá nulový.

Do trvalého stavu se tedy řetězec vrátí s pravděpodobností 1 po konečně mnoha krocích, zatímco do přechodného stavu se s kladnou pravděpodobností nikdy nevrátí. Z věty 3 plyne následující tvrzení.

Důsledek 1. V nerozložitelném aperiodickém Markovově řetězci jsou všechny stavy trvalé nenulové.

Příklad 4. Model havarijního pojištění. Nechť pojišťovna používá tři kategorie pojistného pro pojištění osobních automobilů: s_1 základní pojistné, s_2 bonus 30 %, s_3 bonus 50 %. Průběh platby pojistného daného pojištěnce lze modelovat Markovovým řetězcem s množinou stavů $\{s_1, s_2, s_3\}$, který startuje ze stavu s_1 při uzavření pojistky. Přechody o kategorii výše jsou podmíněny bezškodním průběhem v daném roce, nastane-li aspoň jedna pojistná událost, je pojištěný zařazen v dalším roce do stavu s_1 . Předpokládejme, že počet výskytů pojistných událostí v roce je náhodná veličina s Poissonovým rozdělením pravděpodobnosti s parametrem λ . Potom matice pravděpodobností přechodu má tvar

$$\mathbb{P} = \begin{pmatrix} 1 - e^{-\lambda} & e^{-\lambda} & 0 \\ 1 - e^{-\lambda} & 0 & e^{-\lambda} \\ 1 - e^{-\lambda} & 0 & e^{-\lambda} \end{pmatrix}. \quad (4)$$

řetězec je nerozložitelný aperiodický a všechny stavy jsou trvalé nenulové.

Příklad 5. Markovův řetězec X z příkladu 2 není nerozložitelný. Stav 6 je trvalý, neboť pravděpodobnost, že padne 6, je kladná, a jakmile se řetězec dostane do stavu 6, už v něm setrvá (říkáme též, že 6 je pohlcující stav). Naproti tomu ostatní stavy 1,2,3,4,5 jsou z tohoto důvodu přechodné.

4. Simulace Markovova řetězce

Nechť Markovův řetězec X má přechodovou matici $\mathbb{P}$ a počáteční rozdělení $\xi^{(0)}$. Předpokládejme, že $U_0, U_1, \dots$ je posloupnost nezávislých náhodných veličin s rovnoměrným rozdělením na intervalu $\langle 0, 1 \rangle$, nezávislých též na X_0 . V simulacích je tato posloupnost aproximována generátorem pseudonáhodných čísel. Zabýváme se otázkou, jak simulovat realizaci řetězce X . Výběr z počátečního rozdělení se realizuje pomocí zobrazení $f : \langle 0, 1 \rangle \mapsto S$ takového, že $f(x) = s_1$ pro $x \in \langle 0, \xi_1^{(0)} \rangle$, $f(x) = s_2$ pro $x \in \langle \xi_1^{(0)}, \xi_1^{(0)} + \xi_2^{(0)} \rangle, \dots$

$f(x) = s_k$ pro $x \in \langle \sum_{j=1}^{k-1} \xi_j^{(0)}, 1 \rangle$. Potom položíme $X_0 = f(U_0)$ a platí

$$P(X_0 = s_j) = P(f(U_0) = s_j) =$$

$$E1_{[f(X_0)=s_j]} = \int_0^1 1_{[f(x)=s_j]} dx = \xi^{(0)}(s_j)$$

pro každé $j = 1, \dots, k$.

Přechod od X_n k X_{n+1} pro libovolné $n \geq 0$ se realizuje funkcí $\phi : S \times \langle 0, 1 \rangle \mapsto S$ definované pomocí matice $\mathbb{P}$. Volíme $\phi(s_i, x) = s_1$ pro $x \in \langle 0, p_{i1} \rangle$, $\phi(s_i, x) = s_2$ pro $x \in \langle p_{i1}, p_{i1} + p_{i2} \rangle, \dots$ $\phi(s_i, x) = s_k$ pro $x \in \langle \sum_{j=1}^{k-1} p_{ij}, 1 \rangle$, $i = 1, \dots, k$. Simulace je potom dána předpisem

$$X_{n+1} = \phi(X_n, U_{n+1}) \quad (5)$$

a platí skutečně

$$P(X_{n+1} = s_j \mid X_n = s_i) = P(\phi(X_n, U_{n+1}) = s_j \mid X_n = s_i) =$$

$$= P(\phi(s_i, U_{n+1}) = s_j) = \int_0^1 1_{[\phi(s_i, x)=s_j]} dx = p_{ij}.$$

Příklad 6. V příkladu 1 předpokládejme, že řetězec startuje ve stavu s_2 = slušně, tedy $\mu_0 = (0, 1)$. Položíme tedy počáteční funkci $f(x) = s_2$ pro každé $x \in \langle 0, 1 \rangle$. K simulaci dalšího chodu volíme např.

$$\phi(s_1, x) = \begin{cases} s_1, & x \in \langle 0, \frac{3}{5} \rangle \\ s_2, & x \in \langle \frac{3}{5}, 1 \rangle \end{cases}$$

a

$$\phi(s_2, x) = \begin{cases} s_1, & x \in \langle 0, \frac{3}{10} \rangle \\ s_2, & x \in \langle \frac{3}{10}, 1 \rangle. \end{cases}$$

5. Stacionární rozdělení

Definice 5. Vektor $\pi = (\pi_1, \dots, \pi_k)$ se nazývá stacionární rozdělení Markovova řetězce X , jestliže má nezáporné prvky se $\sum_{i=1}^k \pi_i = 1$ a platí

$$\pi \mathbb{P} = \pi. \quad (6)$$

Příklad 7. V příkladu 1 k dané přechodové matici dostáváme soustavu rovnic (6) ve tvaru

$$\begin{aligned} 0,6\pi_1 + 0,3\pi_2 &= \pi_1 \\ 0,4\pi_1 + 0,7\pi_2 &= \pi_2, \end{aligned}$$

jejímž řešením je stacionární rozdělení ve tvaru $\pi_1 = \frac{3}{7}$, $\pi_2 = \frac{4}{7}$.

Lze ukázat, že stacionární rozdělení vždy existuje. Zabývejme se limitním chováním posloupnosti rozdělení $\xi^{(n)}$. Ukáže se, že za určitých předpokladů tato posloupnost konverguje ke stacionárnímu rozdělení.

Definice 6. Pro dvě pravděpodobnostní rozdělení $\chi = (\chi_1, \dots, \chi_k)$, $\eta = (\eta_1, \dots, \eta_k)$ na S definujeme vzdálenost ve smyslu totální variace d_{TV} jako

$$d_{TV}(\chi, \eta) = \frac{1}{2} \sum_{i=1}^k |\chi_i - \eta_i|. \quad (7)$$

Posloupnost rozdělení $\chi^{(n)}$ konverguje k χ v totální variaci jestliže

$$\lim_{n \rightarrow \infty} d_{TV}(\chi^{(n)}, \chi) = 0.$$

Píšeme $\chi^{(n)} \xrightarrow{TV} \chi$.

Věta 4. Pro libovolné stacionární rozdělení π nerozložitelného aperiodického Markovova řetězce X platí

$$\xi^{(n)} \xrightarrow{TV} \pi. \quad (8)$$

Důkaz (Häggström, 2002): Nechť π je stacionární rozdělení řetězce X , který startuje z rozdělení $\xi^{(0)}$, tj.

$$X_0 = f_{\xi^{(0)}}(U_0), \quad X_1 = \phi(X_0, U_1), \dots$$

pro funkce $f = f_{\xi^{(0)}}$, ϕ a posloupnost $U_0, U_1, \dots$ z popisu simulace Markovova řetězce. Vezměme nezávisle na této posloupnosti ještě druhou takovou posloupnost $U'_0, U'_1, \dots$ a generujme řetězec X' předpisem

$$X'_0 = f_{\pi}(U'_0), \quad X'_1 = \phi(X'_0, U'_1), \dots,$$

tedy s počátečním rozdělením π . X' má potom rozdělení π v libovolném čase n , dále X, X' jsou nezávislé.

Ukážeme, že s pravděpodobností 1 se oba řetězce setkají v konečném čase. Označme $T = \min\{n; X_n = X'_n\}$, $T = \infty$ pokud se nesetkají. Z věty 2 existuje M přirozené tak, že $p_{ij}^M > 0$ pro každé $i, j \in \{1, \dots, k\}$. Označme $\alpha = \min\{p_{ij}^M; i, j \in \{1, \dots, k\}\}$, je $\alpha > 0$. Dostáváme

$$\begin{aligned} P(T \leq M) &\geq P(X_M = X'_M) \geq P(X_M = s_1, X'_M = s_1) = \\ &= P(X_M = s_1)P(X'_M = s_1) = \end{aligned}$$

$$\begin{aligned}
&= \sum_{i=1}^k P(X_0 = s_i, X_M = s_1) \sum_{i=1}^k P(X'_0 = s_i, X'_M = s_1) = \\
&= \sum_{i=1}^k P(X_0 = s_i) P(X_M = s_1 \mid X_0 = s_i) \times \\
&\quad \times \sum_{i=1}^k P(X'_0 = s_i) P(X'_M = s_1 \mid X'_0 = s_i) \\
&\geq \left(\alpha \sum_{i=1}^k P(X_0 = s_i) \right) \left(\alpha \sum_{i=1}^k P(X'_0 = s_i) \right) = \alpha^2,
\end{aligned}$$

takže $Pr(T > M) \leq 1 - \alpha^2$. Analogicky $P(X_{2M} \neq X'_{2M} \mid T > M) \leq 1 - \alpha^2$, tedy

$$\begin{aligned}
P(T > 2M) &= P(T > M) P(T > 2M \mid T > M) \leq \\
(1 - \alpha^2) P(T > 2M \mid T > M) &\leq (1 - \alpha^2) P(X_{2M} \neq X'_{2M} \mid T > M) \leq \\
(1 - \alpha^2)^2. &\text{ Tak dostáváme pro libovolné } l : P(T > lM) \leq (1 - \alpha^2)^l \text{ a konečně}
\end{aligned}$$

$$\lim_{n \rightarrow \infty} P(T > n) = 0,$$

což znamená, že řetězce se s pravděpodobností 1 setkají.

K dokončení důkazu sestrojíme třetí řetězec $X'' = (X''_0, X''_1, \dots)$ tak, že položíme $X''_0 = X_0$ a dále

$$X''_{n+1} = \phi(X''_n, U_{n+1}) \quad \text{pokud } X''_n \neq X'_n,$$

$$X''_{n+1} = \phi(X''_n, U'_{n+1}) \quad \text{pokud } X''_n = X'_n.$$

X'' je Markovův řetězec, neboť U_i, U'_i jsou nezávislé, navíc má rozdělení $\xi^{(n)}$ jako X pro každé n . Necht' $i \in \{1, \dots, k\}$, dostáváme

$$\begin{aligned}
\xi_i^{(n)} - \pi_i &= P(X''_n = s_i) - P(X'_n = s_i) \leq \\
&\leq P(X''_n = s_i, X'_n \neq s_i) \leq P(X''_n \neq X'_n) = P(T > n),
\end{aligned}$$

což konverguje k nule. Analogicky $\pi_i - \xi_i^{(n)} \leq P(T > n)$, a celkem

$$\lim_{n \rightarrow \infty} |\xi_i^{(n)} - \pi_i| = 0.$$

Odtud

$$\lim_{n \rightarrow \infty} d_{TV}(\xi^{(n)}, \pi) = \lim_{n \rightarrow \infty} \left(\frac{1}{2} \sum_{i=1}^k |\xi_i^{(n)} - \pi_i| \right) = 0,$$

a věta je dokázána.

Za předpokladů věty máme tedy bez ohledu na tvar počátečního rozdělení konvergenci $\xi^{(n)}$ ke stacionárnímu rozdělení.

Důsledek 2. *Nerazložitelný aperiodický Markovův řetězec má právě jedno stacionární rozdělení.*

Důkaz: Necht π, π' jsou dvě stacionární rozdělení řetězce X . Jestliže řetězec startuje z $\xi^{(0)} = \pi'$, je $\xi^{(n)} = \pi'$ pro každé n ze stacionarity π' . Z věty 4 ovšem plyne $\xi^{(n)} \xrightarrow{TV} \pi$, tedy

$$\lim_{n \rightarrow \infty} d_{TV}(\pi', \pi) = 0.$$

Odsud $\pi = \pi'$.

Příklad 8. V příkladu 4 o havarijním pojištění najdeme jediné stacionární rozdělení řetězce s maticí pravděpodobností přechodu (4). řešením soustavy rovnic $\pi \mathbb{P} = \pi$ dostáváme pravděpodobnostní rozdělení

$$(\pi_1; \pi_2; \pi_3) = (1 - e^{-\lambda}, e^{-\lambda} - e^{-2\lambda}, e^{-2\lambda}).$$

Parametr λ je charakteristika řidiče; s rostoucím λ se zvětšuje střední počet pojistných událostí. Pro $\lambda = 0,1$ resp. $0,5$ je stacionární rozdělení rovno $(0,0952; 0,0861; 0,8187)$ resp. $(0,3935; 0,2386; 0,3679)$. Tato čísla udávají přibližně podíly pojištěných v jednotlivých kategoriích při dlouhém chodu pojištění.

Příklad 9. Uvažme Markovův řetězec se dvěma stavy s_1 = stroj je v provozu, s_2 = stroj je mimo provoz. Necht pravděpodobnosti změny stavu jsou α, β z $(0, 1)$, matice pravděpodobností přechodu

$$\mathbb{P} = \begin{pmatrix} 1 - \alpha & \alpha \\ \beta & 1 - \beta \end{pmatrix}. \quad (9)$$

Potom stacionární rozdělení vychází (srovnej příklad 7)

$$\pi = \left(\frac{\beta}{\alpha + \beta}, \frac{\alpha}{\alpha + \beta} \right). \quad (10)$$

6. Rychlost konvergence

V mnohých aplikacích pravděpodobnosti a statistiky je zapotřebí simulovat ze složitěho rozdělení pravděpodobnosti (např. mnohorozměrného), kde není možný přímý postup simulace jako výše pro počáteční rozdělení. Potom se používají metody nazývané Markovskými Monte Carlo (anglicky Markov chain Monte Carlo, zkratka MCMC), kde hledané rozdělení je limitní (stacionární) pro nějaký Markovův řetězec, jehož chod simulovat lze. Zde hraje významnou roli rychlost konvergence vyšetřované v předchozím odstavci o stacionárním rozdělení, neboť v praxi je realizovaný chod řetězce vždy konečný. Poznamenejme, že pro malou třídu podobných úloh lze stacionární rozdělení dosáhnout simulací řetězce konečné délky (což se nazývá exaktní simulace, viz Haggström, 2002), ale tyto metody přesahují rámec našeho textu. Pro velkou většinu úloh zůstává studium rychlosti konvergence ke stacionárnímu rozdělení podstatné, proto uvádíme úvod do této problematiky.

V případě konečného stavového prostoru vede problém k vyšetřování vlastních čísel a vlastních vektorů matice pravděpodobností přechodu. Připomeňme, že pokud pro matici A rozměru $r \times r$ existují skalár $\lambda \in \mathbb{C}$ ($\mathbb{C}$ je množina komplexních čísel) a nenulový vektor $v \in \mathbb{C}^r$ takový, že

$$Av^T = \lambda v^T, \quad \text{resp.} \quad vA = \lambda v,$$

nazývá se v pravý resp. levý vlastní vektor A příslušný vlastnímu číslu λ . Zde T značí transpozici řádkového vektoru na sloupcový. Má-li A navzájem různá vlastní čísla, označme $u_1, \dots, u_r$ resp. $v_1, \dots, v_r$ sadu jejích levých resp. pravých vlastních vektorů takových, že $u_i v_i^T = 1$, $i = 1, \dots, r$. Potom máme spektrální rozklad (Gantmacher, 1959)

$$A^n = \sum_{i=1}^r \lambda_i^n v_i^T u_i. \quad (11)$$

Příklad 10. Uvažme matici pravděpodobností přechodu (9) z příkladu 9. Její charakteristický mnohočlen $(1 - \alpha - \lambda)(1 - \beta - \lambda) - \alpha\beta$ dává kořeny $\lambda_1 = 1$ a $\lambda_2 = 1 - \alpha - \beta$. Vlastnímu číslu $\lambda_1 = 1$ odpovídá pravý vlastní vektor $v = \mathbf{1} = (1, 1)$ a levý vlastní vektor (10) (stacionární rozdělení). Rozklad (11) má tvar (viz též Resnick, 1992, str. 76)

$$\mathbb{P}^n = \frac{1}{\alpha + \beta} \begin{pmatrix} \beta & \alpha \\ \beta & \alpha \end{pmatrix} + \frac{(1 - \alpha - \beta)^n}{\alpha + \beta} \begin{pmatrix} \alpha & -\alpha \\ -\beta & \beta \end{pmatrix}$$

a jelikož $|1 - \alpha - \beta| < 1$, je (konvergenci matic uvažujeme po prvcích)

$$\lim_{n \rightarrow \infty} \mathbb{P}^n = \mathbb{P}^\infty = \frac{1}{\alpha + \beta} \begin{pmatrix} \beta & \alpha \\ \beta & \alpha \end{pmatrix}.$$

Tato algebraická metoda dává

$$\mathbb{P}^n - \mathbb{P}^\infty = \frac{(1 - \alpha - \beta)^n}{\alpha + \beta} \begin{pmatrix} \alpha & -\alpha \\ -\beta & \beta \end{pmatrix},$$

tedy

$$\sup_{i,j} |p_{ij}^{(n)} - \pi_j| \leq K(1 - \alpha - \beta)^n,$$

kde $K > 0$ je konstanta. Ukázali jsme, že rychlost konvergence je geometrická a určená modulem druhého největšího vlastního čísla. S uvažováním věty 1 (pro libovolné počáteční rozdělení) tedy geometrická rychlost konvergence platí i v (8), tj. pro absolutní pravděpodobnosti $\xi^{(n)}$ ve větě 4.

Výsledek z příkladu zobecňuje Perron-Frobeniova věta, kterou formuluje speciálně pro matici pravděpodobností přechodu:

Věta 5. *Nechť $\mathbb{P}$ je matice typu $r \times r$, která je maticí pravděpodobností přechodu nerozložitelného aperiodického Markovova řetězce. Potom $\mathbb{P}$ má vlastní číslo $\lambda_1 = 1$ násobnosti jedna a ostatní vlastní čísla $\lambda_2, \dots, \lambda_r$ lze seřadit sestupně $1 = \lambda_1 > |\lambda_2| \geq |\lambda_3| \geq \dots \geq |\lambda_r|$. Platí*

$$\mathbb{P}^n = \mathbf{1}^T \pi + O(n^{m_2-1} |\lambda_2|^n),$$

kde m_2 je násobnost λ_2 a $O(f(n))$ je funkce n taková, že existují $\alpha, \beta \in \mathbb{R}$ s $0 < \alpha \leq \beta < \infty$, pro něž platí $\alpha f(n) \leq O(f(n)) \leq \beta f(n)$ pro všechna dostatečně velká n .

Důkaz: Gantmacher (1959).

U řetězců s rozsáhlou množinou stavů a složitou přechodovou strukturou nemusí být určení λ_2 přímočaré, potom se rychlost konvergence určuje různými odhady, viz Brémaud (1998, kapitola 6).

Literatura

- [1] Anděl J. (1976). *Statistická analýza časových řad*. SNTL, Praha.
- [2] Brémaud P. (1998). *Markov Chains: Gibbs fields, Monte Carlo Simulation, and Queues*. Springer, New York.
- [3] Gantmacher F.R. (1959). *Applications of the Theory of Matrices*. Chelsea, New York (překlad z ruštiny).
- [4] Häggström O. (2003). *Finite Markov Chains and Algorithmic Applications*. Cambridge Univ. Press, Cambridge.
- [5] Resnick S. (1992). *Adventures in Stochastic Processes*. Birkhäuser, Basel.

ZÁKLADY STATISTICKÉHO MYŠLENÍ

Daniel Hlubinka

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

daniel.hlubinka@mff.cuni.cz

1. Náhodný výběr

Ve statistice je ústředním pojmem náhodný výběr.

Definice 1. Náhodným výběrem rozumíme posloupnost (vektor)

$$\mathbf{X} = (X_1, X_2, \dots, X_n) \quad (1)$$

nezávislých náhodných veličin se stejným rozdělením. Číslo n se nazývá *rozsah* náhodného výběru.

Poznamenejme, že takové náhodné veličiny se označují též zkratkou *iid* (z anglického *independent and identically distributed*). Náhodný výběr je matematickým modelem pokusů prováděných opakovaně a nezávisle za stejných podmínek. V praxi pak dostáváme

$$\mathbf{x} = (x_1, x_2, \dots, x_n) \quad (2)$$

realizaci náhodného výběru. Jinými slovy naměřené hodnoty x_i náhodných veličin X_i .

Zde můžeme poznamenat, že náhodná veličina je takto „dvojjediná“, což mnohdy působí potíže při studiu statistiky. Na jednu stranu zde máme funkci $X : \Omega \rightarrow \mathbb{R}$ přiřazující elementárním jevům teoretického pravděpodobnostního prostoru nějakou reálnou hodnotu. Tedy X je v tomto pojetí funkcí. Na druhou stranu jako funkci náhodnou veličinu nikdy nepozorujeme. Vždy máme pozorování $X = x$, tedy reálnou hodnotu (v případě náhodného výběru vektor hodnot) a náhodnou veličinu vnímáme často právě jako hodnotu, byť předem nevíme, jaká hodnota bude pozorována.

Klíčová slova: Náhodný výběr, bodový odhad, test statistické hypotézy.

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839 a grantu GAČR 201/03/p138.

Náhodný výběr reprezentuje nějakou teoretickou náhodnou veličinu s jejím rozdělením. To znamená, že reprezentuje rozdělení této veličiny, je pozorovaným uskutečněním náhody. Proto se též často říká *náhodný výběr z ... rozdělení*, kde slovu rozdělení předchází specifikace pravděpodobnostního modelu. Můžeme tomu také rozumět tak, že náhodný výběr je opakovanou replikou náhodné veličiny a každá realizace náhodné veličiny přispívá k informaci o jejím rozdělení. Pro náhodnou veličinu s nějakým specifikovaným rozdělením budeme používat značení $X \sim \dots$, jako zkratku za výrok „náhodná veličina X má rozdělení“ ...

Příklad 1. Při kontrole kvality výrobku je prováděna kontrola u n náhodně vybraných výrobků. Výsledkem kontroly je zjištění, zda výrobek je bez chyby (bezvadný) či nikoliv. Matematizovat takový proces je možné například takto

$$X_i = \begin{cases} 0 & \text{je-li } i\text{-tý výrobek bezvadný,} \\ 1 & \text{je-li } i\text{-tý výrobek kazový.} \end{cases}$$

Realizace náhodného výběru je tedy prvek množiny $\{0, 1\}^n$. Navíc zpravidla předpokládáme, že pravděpodobnost výskytu kazů u výrobku je stálá, a to $p = P(X_1 = 1)$. Náhodné veličiny X_i tedy mají alternativní rozdělení, a kazovosti jednotlivých výrobků jsou nezávislé jevy. Máme tedy výběr z alternativního rozdělení, neboli, pozorujeme alternativně rozdělenou náhodnou veličinu.

K čemu potřebujeme nezávislost? Vztahy (závislost) mezi jednotlivými náhodnými veličinami mohou nabývat nepřeborného množství podob. Je však jeden výlučný vztah, jenž umožňuje snadno určit sdružené rozdělení náhodných veličin a tím je právě nezávislost. Připomeňme si, že pravděpodobnost průniku nezávislých jevů je rovna součinu pravděpodobností těchto jevů.

Příklad 1 (pokračování). Realizací kontroly n výrobků je n -rozměrný vektor $(x_1, \dots, x_n)$ sestávající se z nul a jedniček. Pravděpodobnost jedné takové realizace s právě k jedničkami (kazovými výrobky) na daných místech je rovna díky předpokladu nezávislosti $p^k(1-p)^{n-k}$, tj. součinu pravděpodobností výskytu kazu (k -krát) a pravděpodobností bezvadného výrobku $((n-k)$ -krát). Nás však spíše než pravděpodobnost konkrétní realizace zajímá pravděpodobnost, že při kontrole dojde k odhalení daného počtu kazů. Označme proto Y celkový počet kazů ve výběru, tedy $Y = \sum_{i=1}^n X_i$. Snadno zjistíme, že pravděpodobnost, že mezi n výrobky bude právě k vadných je

$$P(Y = k) = \binom{n}{k} p^k (1-p)^{n-k}, \quad k \in \{0, 1, \dots, n\}. \quad (3)$$

Rozdělení Y je tedy nám již známé binomické rozdělení $B(n, p)$.

Viděli jsme, že nezávislost nám umožnila přesně spočítat pravděpodobnost výskytu daného počtu zmetků mezi kontrolovanými výrobky. Podobně je nezávislost užitečná i v jiných případech. Můžeme tedy shrnout, že nezávislost je tím přirozeným prostředkem, který nám umožní snadněji provádět statistickou analýzu náhodného výběru. Díky znalosti rozdělení můžeme pak například navrhnout metodu pro odhad parametru pravděpodobnostního modelu, nebo o něm testovat hypotézu. Nezávislost je navíc rozumným a očekávaným předpokladem pro chování posloupnosti pokusů začínajících vždy „od začátku“.

2. Problém parametru

Předpokládejme, že máme dobrý důvod předpokládat určitý pravděpodobnostní model, tedy rozdělení náhodné veličiny, a že jsme schopni opakovaně, nezávisle a za stejných podmínek opakovat náhodný pokus, a tak vytvořit náhodný výběr. Již v příkladu 1 si můžeme povšimnout, že zde stále zůstává jedna neznámá již je hodnota $p \in (0, 1)$. V obecné rovině pak můžeme mluvit o parametrické statistice. V předpokládaném modelu rozdělení náhodné veličiny zůstává neznámý parametr o němž chceme na základě náhodného výběru zjistit co nejvíce.

Definice 2. Parametrickou třídou rozdělení náhodné veličiny X rozumíme systém distribučních funkcí

$$F_{\theta}(x) = P_{\theta}(X \leq x), \quad \theta \in \Theta, \quad (4)$$

kde hodnotu θ nazýváme parametrem, množinu či prostor Θ parametrickým prostorem (nejčastěji podmnožina d -rozměrného Euklidova prostoru a d je nějaké (většinou malé) přirozené číslo). O parametru θ předpokládáme, že je pevný (jednotlivá pozorování náhodného výběru mají stejné rozdělení), avšak není přímo pozorovatelný. Stejně tak není pozorovatelná hypotetická *populace*, kterou zde rozumíme soubor všech možných realizací náhodné veličiny s daným rozdělením.

Poznamenejme, že statistika v tomto smyslu je (matematizovanou) metodou induktivního myšlení. Z pozorovaných realizací náhodné veličiny (náhodného výběru) usuzujeme na neznámé hypotetické vlastnosti populace (rozdělení náhodné veličiny). V příkladu 1 by nás pochopitelně zajímalo, jaký je podíl kazových výrobků, případně rozhodnutí, zda tento podíl překračuje nějakou předem zvolenou mez. O parametru p si však můžeme udělat představu

jen na základě počtu napozorovaných vadných výrobků mezi n kontrolovanými, jiné údaje nemáme.

K odhadu parametru z náhodného výběru nám mohou napomoci charakteristiky náhodných veličin známé z předchozích kapitol. Zatímco u pravděpodobnostního modelu se na první pohled (samozřejmě nespravedlivě) mohou jevit jako poněkud nadbytečné, neboť úplná znalost rozdělení je dostačující, nyní na porovnání populačních a výběrových charakteristik postavíme jednu z metod odhadování parametrů rozdělení.

Definice 3. Výběrovým průměrem pro náhodný výběr $\mathbf{X} = (X_1, \dots, X_n)$ rozumíme hodnotu

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i, \quad (5)$$

tedy aritmetický průměr pozorování.

Výběrovým rozptylem rozumíme hodnotu

$$M_2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2, \quad (6)$$

a nestranným výběrovým rozptylem pak

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \frac{n}{n-1} M_2, \quad (7)$$

pokud $n \geq 2$, jinak tato charakteristika nemá smysl (což ovšem v tom případě nemá ani $M_2 \equiv 0$). Podobně lze definovat další centrální výběrové momenty pro $n \in \mathbb{N}$.

$$M_n = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^n \quad (8)$$

Poznamenejme, že výběrový průměr a rozptyl jsou opět náhodné veličiny a také jejich rozdělení závisí na parametru θ . Obecně však toto rozdělení nemusí být ani ze stejné třídy, jako rozdělení sledované náhodné veličiny. Zopakujme, co už víme.

Lemma 1. *Nechť $X_1, X_2, \dots, X_n$, $n \geq 2$ je náhodný výběr a dále necht' $EX_1 = \mu$ a $\text{var}X_1 = \sigma^2$. Pak platí*

$$E\bar{X} = \mu, \quad \text{var}\bar{X} = \frac{\sigma^2}{n}, \quad ES^2 = \frac{n}{n-1} EM_2 = \sigma^2 \quad (9)$$

Z lemmatu 1 vidíme, proč se S^2 nazývá nestranný výběrový rozptyl.

Poznamenejme zde, že výběrový průměr a oba výběrové rozptyly z definice 3 jsou v tomto pojetí náhodné veličiny. Pro pozorovanou realizaci náhodného výběru pak pochopitelně jde o reálná čísla, ne náhodné veličiny. Tato pozorování označíme malými písmeny

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i, \quad m_2 = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2, \quad s^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2.$$

I zde můžeme poznamenat, že výběrové momenty jsou „dvojjediné“ ve stejném smyslu jako náhodná veličina, kterou ve skutečnosti jsou.

Střední hodnota a rozptyl náhodné veličiny bývají funkcí parametru rozdělení, $\mu = \mu_1(\theta)$, $\sigma^2 = \mu_2(\theta)$. Tato skutečnost je základem jednoduchého způsobu odhadu $\hat{\theta}$ parametru θ , takzvané *momentové metody*.

3. Bodový odhad parametru

Definice 4. Mějme náhodný výběr z parametrické třídy F_θ . Odhadem $\hat{\theta}$ parametru θ momentovou metodou je hodnota splňující

$$\bar{X} = \mu_1(\hat{\theta}), \quad \hat{\theta} \in \Theta. \quad (10)$$

Je-li parametr θ vícerozměrný a (10) nestačí pro jednoznačné určení odhadu, pokračujeme rovnostmi

$$M_k = \mu_k(\hat{\theta}), \quad k = 2, 3, \dots \quad (11)$$

dokud se nám nepodaří určit $\hat{\theta}$ jednoznačně.

Poznamenejme, že obvykle je θ prvkem podmnožiny reálných čísel a μ_1 je jednoznačné zobrazení a proto je použití momentové metody v těchto případech snadné. I v obecném případě, kdy je parametr vícesložkový, jako například u normálního rozdělení, stačí použít několik momentových rovností a získáme hledaný odhad.

Příklad 1 (pokračování). Vraťme se ke kontrole kvality výrobků. Již z minula víme, že pro $X \sim \text{Alt}(p)$ platí $EX = p$. Proto snadno vidíme, že $\hat{p} = \bar{X}$.

Hned ale vidíme také další vlastnosti. Odhad je nestranný, neboť $E\hat{p} = p$ pro všechny hodnoty parametru. Platí také $\text{var } \hat{p} = p(1-p)/n$, neboli pro každé p rozptyl odhadu konverguje k nule. Z poučení v kapitole o základech pravděpodobnosti lze ovšem vyvodit více: odhad $\hat{p}$ je náhodná veličina založená na náhodném výběru o rozsahu n a při zvětšujícím se rozsahu výběru

tato veličina konverguje ke skutečné hodnotě parametru v pravděpodobnostním smyslu. Tomu se říká *konzistence odhadu*.

Dodejme hned, že obecně odhad momentovou metodou nemusí být nestranný ani konzistentní. Kdybychom chtěli odhadnout rozptyl $p(1-p)$, můžeme postupovat dvěma způsoby. Buď použít výběrový rozptyl M_2 , nebo odhad založený na $\hat{p}(1-\hat{p}) = \bar{X}(1-\bar{X})$. V prvním případě dostáváme

$$\begin{aligned} EM_2 &= \frac{1}{n} \sum_{i=1}^n [EX_i^2 - 2EX_i\bar{X} + E\bar{X}^2] = EX_1^2 - 2EX_1\bar{X} + E\bar{X}^2 \\ &= p - \frac{2}{n} \sum_{i=1}^n EX_1X_i + \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n EX_iX_j \\ &= p - \frac{2}{n}p - \frac{2(n-1)}{n}p^2 + \frac{1}{n} \sum_{j=1}^n EX_1X_j \\ &= p - \frac{2}{n}p - \frac{2(n-1)}{n}p^2 + \frac{1}{n}p + \frac{(n-1)}{n}p^2 = \frac{n-1}{n}p(1-p). \end{aligned} \quad (12)$$

Stejný výsledek, tedy $EM_2 = E\bar{X}(1-\bar{X})$, obdržíme i v druhém případě. Ve skutečnosti jsou v tomto případě oba odhady shodné, neboť v alternativním rozdělení platí $P[X=1] = P[X^2=1]$ a $P[X=0] = P[X^2=0]$. Jednotlivé kroky výpočtu (12) využívají linearitu střední hodnoty a nezávislost jednotlivých proměnných. Dostáváme tedy, že odhad rozptylu provedený momentovou metodou není nestranný, na druhou stranu vidíme přesně, co mu do nestrannosti „chybí“. Platí totiž

$$EM_2 = \frac{n-1}{n}p(1-p) = \text{var}X - \frac{1}{n}p(1-p). \quad (13)$$

Člen $p(1-p)/n$, odchylka od požadované hodnoty, se nazývá *vychýlením odhadu*. Známe-li vychýlení, není těžké odhad upravit na nestranný, zde vynásobením konstantou $n/(n-1)$. Je však otázkou, zda je zde oprava nutná. Všimněme si totiž, že pro všechna n je vychýlení menší než $(4n)^{-1}$, tedy pro velké rozsahy výběru začíná být zanedbatelné. Ať už s opravou, nebo bez ní, náš odhad rozptylu je konzistentní. Konverguje-li $\bar{X}$ v pravděpodobnosti k p , pak i každá spojitá funkce (tedy i $\bar{X}(1-\bar{X})$) konverguje k odpovídající hodnotě $p(1-p)$. Teorie odhadu parametrů je pěkně vysvětlena i s odpovídajícím matematickým aparátem například v knihách [1] a [2].

Rozlišujme způsob provedení odhadu a jeho vlastnosti. Ukázali jsme si momentovou metodu jako jednoduchý (nejde však o jediný možný) způsob jak

odhad provést. Získaný odhad může mít různé kvality. Může jít o nestrannost, kdy střední hodnota odhadu jako náhodné veličiny odpovídá odhadované neznámé hodnotě. Jinou užitečnou vlastností odhadu je jeho konzistence, kdy se zvětšujícím se rozsahem n náhodného výběru konverguje odhad k odhadované hodnotě (tedy jako náhodná veličina degeneruje do bodu). Poznamenejme, že odhady stejného parametru získané různými metodami mohou mít stejné kvality. Na druhou stranu mohou mít odhady parametru v různých modelech, byť provedené stejnou metodou, různé kvality.

Teď již umíme odhadnout parametr bodově, neznáme ale ještě odpověď na složitější otázky. Typicky se odhad liší od hodnoty, kterou pro tento parametr předpokládáme. Nelze například čekat, že při n hodech kostkou bude zastoupení všech výsledků $n/6$, a to ani v případě, kdy rozsah pokusů je násobkem šesti. Podle zákona velkých čísel však platí, že podíl výsledků se bude blížit ke skutečné hodnotě, pro symetrickou kostku tedy k hodnotě $1/6$. Je-li naše domněnka (hypotéza), že zkoumáme symetrickou kostku, pravdivá, pozorované podíly výsledků se budou limitně blížit k $1/6$. Při konečném počtu pokusů ovšem nezbytně vyvstává otázka na možnou odchylku od teoretického podílu. Jaké odchylky od hypotetické hodnoty ještě „tolerujeme“ jako náhodné, ne v rozporu s hypotézou, a jaké by již byly příliš málo „věrohodně vysvětlitelné“ hypotézou? Tím se dostáváme k problematice testování statistických hypotéz.

4. Testy hypotéz

Začneme příkladem.

Příklad 2. Každý už si asi „házel korunou“ ve chvíli, kdy se nedokázal rozhodnout o dvou v zásadě rovnocenných možnostech. Náhodnost takto učiněné volby je zřejmá, předem nevíme co padne. Předpokládejme, že jsme schopni zaručit při opakovaných pokusech jejich nezávislost. Můžeme si položit otázku, zda mince, kterou házeme je „spravedlivá“, zda oba výsledky, rub a líc, mají stejnou pravděpodobnost, tedy $1/2$? Tedy zda při takto náhodném rozhodování mají obě možnosti stejnou šanci na uskutečnění.

Připomeňme si, že jako pravděpodobnostní model zde předpokládáme alternativní rozdělení s nějakým parametrem $p \in (0, 1)$. My navíc předpokládáme, že ve skutečnosti je $p = 1/2$, dokážeme uskutečnit posloupnost nezávislých n pokusů a z nich odhadnout, dle četnosti předem zvoleného výsledku „líc“, hodnotu p pomocí $\hat{p}$. Parametrický prostor Θ se nám tak rozděluje na část $\Theta_0 = \{1/2\}$ a Θ_1 obsahující zbylé možné hodnoty p . Takovému dělení parametrického prostoru na dvě podmnožiny v rámci modelu říkáme hypotéza

s alternativou.

Definice 5. Necht jsou dány disjunktní množiny $\Theta_0 \subset \Theta$ a $\Theta_1 \subset \Theta$. Uvažujme pouze takový rozklad, že $\Theta_0 \cup \Theta_1 = \Theta$, tedy alternativa je doplňkem hypotézy ve zvoleném modelu. Výrok $\theta \in \Theta_0$ o neznámém parametru $\theta \in \Theta$ nazveme (*nulovou*) *hypotézou* a označíme H_0 , zatímco výrok $\theta \in \Theta_1$ označíme H_1 a nazveme *alternativou*.

Na základě náhodného výběru chceme testovat, zda zvolená nulová hypotéza ob stojí při snaze vysvětlit naměřené hodnoty touto hypotézou. Pokud ano, pak řekneme, že nulovou hypotézu *nezamítáme*. Naopak, pokud hypotéza neobstojí, pak ji *zamítáme*. Rozhodnout se však musíme, nepřipouštíme nějaké poloviční či podmíněné zamítnutí (nezamítnutí) hypotézy. Je tedy zřejmé již z vybraných pojmů, že hypotézu zamítneme tehdy, máme-li k tomu dostatečně silné argumenty. Jinak, poněkud opatrně, říkáme, že nemáme dostatek důkazů o neplatnosti H_0 . S trochou nadsázky můžeme říci, že pro nulovou hypotézu platí presumpce nevinu.

Je zřejmé, že při takovém dichotomickém rozhodování celkem mohou nastat čtyři možnosti.

	H_0 platí	H_0 neplatí
H_0 nezamítáme	správné rozhodnutí	špatné rozhodnutí
H_0 zamítáme	špatné rozhodnutí	správné rozhodnutí

Zatímco u výroků o platnosti hypotézy nevíme, zda jsou pravdivé, jsme to my, kdo rozhoduje o zamítnutí či nezamítnutí hypotézy. Tabulka ukazuje, kdy je naše rozhodnutí v pořádku a proto bychom chtěli, aby takových rozhodnutí bylo co nejvíce. Naopak chybná rozhodnutí a jejich výskyt chceme co nejvíce omezit. Jak jsme již řekli výše, nejsou zde hypotéza a její alternativa v rovnocenném postavení a proto i obě chyby mají různé postavení ve statistice.

Protože naše rozhodování je postaveno na náhodném výběru, je nutné o něm uvažovat opět jako o náhodném jevu. Testování hypotéz je v tomto pohledu optimalizační úlohou, kdy chceme maximalizovat pravděpodobnost správného rozhodnutí a minimalizovat pravděpodobnost špatného rozhodnutí. Ukazuje se, že je možné (a mnohdy i optimální) rozdělit prostor všech možných výsledků $\mathcal{W}$ náhodného výběru $\mathbf{X}$ na dvě disjunktní části $\mathcal{W}_0$ a $\mathcal{W}_1$, a zároveň tak, aby $\mathcal{W}_0 \cup \mathcal{W}_1 = \mathcal{W}$. Jejich význam je následující. *Padne-li realizace náhodného výběru $\mathbf{x}$ do $\mathcal{W}_0$, nezamítneme H_0 , zatímco při $\mathbf{x} \in \mathcal{W}_1$ zamítáme H_0 a přijímáme H_1 . Množina $\mathcal{W}_1$ se nazývá kritický obor testu.*

Označme $P_\theta(\mathbf{X} \in A)$ pravděpodobnost jevu, že náhodný výběr padne do množiny A pokud je skutečná hodnota parametru θ .

Definice 6. Zamítneme-li platnou hypotézu, dopouštíme se *chyby prvního druhu*. Maximální pravděpodobnost takového rozhodnutí je

$$\alpha = \sup_{\theta \in \Theta_0} P_\theta(\mathbf{X} \in \mathcal{W}_1) \quad (14)$$

a nazývá se *hladinou testu*. Chybou druhého druhu je nezamítnutí neplatné hypotézy. Silofunkcí testu rozumíme

$$\beta(\theta) = P_\theta(\mathbf{X} \in \mathcal{W}_1). \quad (15)$$

Pravděpodobnost chyby druhého druhu se pro $\theta \in \Theta_1$ řídí doplňkem k silofunkci

$$1 - \beta(\theta) = P_\theta(\mathbf{X} \in \mathcal{W}_0), \quad \theta \in \Theta_1. \quad (16)$$

Naším cílem je pochopitelně najít kritický obor testu $\mathcal{W}_1$ tak, aby hladina α byla co nejmenší a síla $\beta(\theta)$ naopak co největší pro všechny hodnoty $\theta \in \Theta_1$. Hned na první pohled je zřejmé, že současně snižovat hladinu a zvětšovat sílu nelze. Nízká hladina testu si vynucuje zmenšování kritického oboru (musíme být opatrní při zamítání hypotézy), což má ale za následek menší sílu. S ohledem na rozdílné postavení hypotézy a alternativy pak postupujeme tak, že si zvolíme hladinu testu a kritický obor hledáme tak, aby síla byla co největší.

Co je vlastně hladina testu? Jde o vědomě podstupované riziko, že zamítneme platnou hypotézu, přičemž míru tohoto rizika hlídáme. Jak ale máme hladinu volit? Můžeme si říci, že jde o sázku, kde zamítnutí platné hypotézy je prohrou, zatímco její nezamítnutí výhrou. Zhodnocení závažnosti této prohry nám pomůže spolu se ziskem z výhry nastavit hladinu testu. Mohu-li vyhrát 10 000 Kč a vsázím jen dnešní oběd, pak přijmu klidně velké riziko prohry. Při stejné výhře nepřistoupím na téměř žádné riziko hrozí-li mi jako prohra měsíční hladovka.

Zcela obecně pro daný model nelze najít test, který by byl optimální pro každou nulovou hypotézu (byť i jen každou jednobodovou). V mnoha základních případech však takový test existuje a jeho provedení je velmi intuitivní.

Příklad 2 (pokračování). Nechť víme, že předložená mince může být buď symetrická ($p = 1/2$), nebo falešná ($p = 1/3$). Chceme zjistit, zda při výsledku tři líce z devíti pokusů můžeme minci prohlásit za falešnou při hladině testu 10 %.

Jak vybrat, co je hypotézou a co alternativou? Pomozme si opět nerovnováhou mezi H_0 a H_1 . Protože hypotézu „chráníme“, její zamítnutí má vyšší vypovídací hodnotu než nezamítnutí. Někteří autoři dokonce nezamítnutí hypotézy vykládají jako „mlčení testu“. Proto tvrzení, které bychom rádi prokázali (i když raději píšme „prokázali“) vždy volíme za alternativu.

Příklad 2 (pokračování). Jest tedy náš test testem

$$H_0 : p = \frac{1}{2} \text{ versus } H_1 : p = \frac{1}{3}. \quad (17)$$

Platnost hypotézy vede k většímu počtu líců spíš nežli platnost alternativy. Proto kritický obor bude tvořen malým počtem líců. Označme X náhodnou veličinu zachycující počet líců v devíti nezávislých pokusech. Pochopitelně $X \sim B(9, p)$, modelem je binomické rozdělení s neznámým parametrem p . Kolik je malý počet líců vedoucí k zamítnutí hypotézy závisí na hladině testu. Kritický obor je pak

$$\mathcal{W}_1 = \{0, 1, \dots, k\}, \text{ kde } k = \max\{x \in \mathbb{N} : P_{H_0}[X \leq x] \leq \alpha\}. \quad (18)$$

Zde tedy je hodnota k rovna dvěma, protože

$$P_{H_0}[X \leq 2] = \left(\frac{1}{2}\right)^9 [1 + 9 + 36] = 0,0898, \quad (19)$$

zatímco $P_{H_0}[X \leq 3] = 0,254$ a tedy bychom překročili hladinu 10 %.

Naše rozhodnutí tedy je *nezamítnout na hladině 10 % hypotézu o symetrické minci*.

Všimněme si, že při volbě libovolné hladiny menší než 8,98 % bychom již kritický obor pozměnili a zamítali bychom hypotézu jen při nejvýše jednom líci z devíti pokusů. Při malém počtu pozorování se může dokonce stát, že hypotézu vůbec zamítnou na dané hladině nelze a to při jakémkoliv výsledku, což je důsledek malého počtu pokusů. Tehdy existuje dostatečně velká pravděpodobnost výskytu každého výsledku pokusu za hypotézy. Například při třech hodech mincí nelze na žádné rozumné hladině (menší než 10 %) zamítnout hypotézu z příkladu 2.

Nicméně je teď na místě **důležité upozornění**. Vzhledem k možnosti ovlivnit rozhodnutí testu je nutné vždy volit hladinu *dříve než provedeme jakékoliv výpočty*. Upravování hladiny tak, aby výsledek testu odpovídal našim představám, je *nepřípustné a neetické*. Podobně by hypotéza a alternativa měly být specifikovány dříve, než začneme daná data analyzovat.

Jednobodová hypotéza a alternativa jsou nejjednodušší možností. Mnohem zajímavější je připuštění více hodnot, zejména pro alternativu.

Příklad 1 (pokračování). Chtěli bychom se ujistit, že podíl zmetků je menší než 0,01. Testovat hypotézu máme na hladině 2 % a z kontroly jsme zjistili, že z 1 000 prověřených výrobků je pět zmetků. Uvažujeme tedy hypotézu a alternativu

$$H_0 : p \geq \frac{1}{100} \text{ versus } H_1 : p < \frac{1}{100}. \quad (20)$$

Proti hypotéze svědčí malé množství zmetků. Proto bude mít kritický obor opět tvar

$$\mathcal{W}_1 = \{0, 1, \dots, k\}, \quad k = \max\{x \in \mathbb{N} : \sup P_{H_0}[X \leq x] \leq 0,02\}. \quad (21)$$

Naštěstí je uvedená pravděpodobnost monotónní v p a největší hodnoty pro hypotézu nabývá při $p = 0,01$. Stačí tedy určit

$$\mathcal{W}_1 = \{0, 1, \dots, k\}, \quad k = \max\{x \in \mathbb{N} : P_{p=0,01}[X \leq x] \leq 0,02\}. \quad (22)$$

Přibližné hodnoty pravděpodobností jednotlivých malých počtů zmetků při $p = 0,01$ jsou

k	0	1	2	3	4
$P[X = k]$	$4 \cdot 10^{-5}$	$4 \cdot 10^{-4}$	0,0022	0,0074	0,0186
$P[X \leq k]$	$4 \cdot 10^{-5}$	$4,4 \cdot 10^{-4}$	0,0026	0,0090	0,0276

(23)

a vidíme, že $P_{p=0,01}[X \leq 5] > 0,02$ a tedy hypotézu o podílu zmetků nejmeně 0,01 zamítnout na hladině 2 % nemůžeme.

Vzhledem k tomu, že jsme si vybrali poměrně rozumná čísla pro náš příklad, mohli jsme počítat výrazy typu

$$\binom{1\,000}{k} 0,01^k 0,99^{1\,000-k} \quad (24)$$

a několik jich sečíst. Přesto je však již na místě uvažovat o zjednodušení výpočtu. Pro binomické rozdělení s velkým počtem pokusů (zde 1 000) a malou pravděpodobností úspěchu (zde 0,01) je dobré použít Poissonovu větu. Pak namísto používání vysokých binomických čísel a počítání velkých mocnin čísel blízkých nule či jedné (což je většinou ošidné), vystačíme si s exponenciálou a mnohem menšími faktoriály a mocninami. Následně mnohem snadněji získáme hodnotu $P_{p=0,01}[X \leq 5] \doteq 0,067$.

Ukázali jsme si takzvanou jednostrannou hypotézu proti jednostranné alternativě. Někdy potřebujeme testovat bodovou hypotézu (předpokládáme jen jednu hodnotu v hypotéze) proti oboustranné alternativě.

Příklad 3. Máme zjistit, zda předložená kostka má pravděpodobnost šestky opravdu $1/6$. Hypotéza tedy zní

$$H_0 : p = \frac{1}{6} \text{ versus } H_1 : p \neq \frac{1}{6}. \quad (25)$$

Všimněme si rozdílu proti předchozímu příkladu. Tam jsme chtěli prokázat, že podíl zmetků je malý a proto jsme tento výrok dali do alternativy. To zde však nejde! Hypotézu $p \neq \frac{1}{6}$ bychom nikdy nemohli zamítnout (proč?). Musíme se tedy spokojit s výsledkem, že kostka, jejíž falešnost není prokázána, může být prozatím považována za spravedlivou, neprokáže-li se časem jinak.

Předpokládejme, že jsme si stanovili hladinu 5 % a provedli 240 pokusů v nichž padlo 48 šestek. Jak se postavit k hypotéze? Proti hypotéze svědčí jak malé, tak i velké množství šestek. Nicméně počítat mnoho malých pravděpodobností založených na hodnotách typu

$$P_{p=1/6}[X = 30] = \binom{240}{30} \frac{5^{210}}{6^{240}} \quad (26)$$

by nám dalo spoustu práce a navíc bychom jistě byli obětí numerické nepřesnosti výpočtu. Mnohem lepší je použití Moivroy-Laplaceovy věty a aproximace rozdělení binomické veličiny Gaussovým rozdělením. V našem případě

$$\begin{aligned} P_{p=1/6}[X \leq x] &= P_{p=1/6} \left[\frac{X - np}{\sqrt{np(1-p)}} \leq \frac{x - np}{\sqrt{np(1-p)}} \right] \\ &= P_{p=1/6} \left[\sqrt{3} \frac{X - 40}{10} \leq \frac{x - 40}{10} \right] \doteq \Phi \left(\frac{x - 40}{10} \right), \end{aligned} \quad (27)$$

kde Φ je distribuční funkce normálního rozdělení. Nyní ovšem víme, že proti hypotéze svědčí jak malé, tak i velké počty šestek, proto kritický obor má tvar

$$\mathcal{W}_1 = \{0, 1, \dots, 40 - k_1\} \cup \{40 + k_2, \dots, 240\}. \quad (28)$$

Díky symetrii normálního rozdělení je vhodné volit $k_1 = k_2 = k$. Určení hodnoty k je snadné, musí platit

$$P_{p=1/6}[X \leq 40 - k] + P_{p=1/6}[X \geq 40 + k] \leq 0,05, \quad (29)$$

příčemž k musí být nejmenší hodnota s takovou vlastností a tedy

$$P_{p=1/6} \left[\sqrt{3} \frac{X-40}{10} \leq \frac{-k\sqrt{3}}{10} \right] + P_{p=1/6} \left[\sqrt{3} \frac{X-40}{10} \geq \frac{k\sqrt{3}}{10} \right] \leq 0,05. \quad (30)$$

Proto hledáme k tak, aby

$$\Phi \left(\frac{-k\sqrt{3}}{10} \right) + 1 - \Phi \left(\frac{k\sqrt{3}}{10} \right) = 2\Phi \left(\frac{-k\sqrt{3}}{10} \right) \leq 0,05. \quad (31)$$

Z tabulek snadno zjistíme, že $\Phi(-1,96) = 0,025$ a tedy $k = 12 \geq 1,96 \cdot 10/\sqrt{3} = 11,31$. Hypotézu bychom tedy zamítli, pokud by počet šestek ve 240 pokusech byl nejvýše 28 nebo nejméně 52. V našem případě na hladině 5 % hypotézu nezamítneme.

5. Vliv velikosti výběru a volby hladiny na test hypotézy

Pojďme si stručně ukázat, jak velikost výběru a zvolená hladina testu ovlivní výsledek testu hypotézy. Pokračujme v příkladu 3.

Příklad 3 (pokračování). Viděli jsme, že na hladině 5 % jsme při 48 šestkách z 240 pokusů nemohli zamítnout hypotézu o spravedlivé kostce. Jak se změní náš závěr při změně hladiny na 10 %? Ověřme-li si v tabulkách, že $\Phi(-1,645) = 0,05$, není již těžké dojít k závěru, že pro $k = 10 \geq 1,64 \cdot 10/\sqrt{3} = 9,47$ platí závěry předchozího příkladu a tedy hypotézu bychom na hladině 10 % zamítali při nejvýše 30 či nejméně 50 šestkách. Tedy do kritického oboru teď patří více výsledků, což odpovídá zvýšené hladině testu, tedy vyšší povolené pravděpodobnosti chyby prvního druhu. Nicméně ani nyní bychom hypotézu o spravedlivé kostce nezamítli.

Vraťme se k hladině 5 % a pokračujme v pokusech. Předpokládejme, že provedeme dalších 120 pokusů a v nich se šestka vyskytne 24 krát. Celkem tedy máme 360 pokusů a v nich 72 šestek. Všimněme si, že podíl šestek se nezměnil, je $1/5$. V testu, provedeném stejným způsobem jako výše, se změní počet pokusů a očekávaný počet šestek při platné hypotéze nyní je 60. Hledáme k tak, aby

$$P_{p=1/6} \left[\frac{X-60}{5\sqrt{2}} \leq \frac{-k}{5\sqrt{2}} \right] + P_{p=1/6} \left[\frac{X-60}{5\sqrt{2}} \geq \frac{k}{5\sqrt{2}} \right] \leq 0,05 \quad (32)$$

a tedy $k = 14 \geq 1,96 \cdot 5 \cdot \sqrt{2} = 13,86$. Hypotézu bychom na hladině 5 % zamítli při nejvýše 46 nebo nejméně 74 šestkách. Všimněme si však, že zatímco rozsah

výběru se zvýšil o polovinu, šířka intervalu, v němž nezamítáme hypotézu vzrostla z 22 na 26, což přibližně odpovídá růstu o rychlosti odmocniny z růstu rozsahu výběru. Tak to je také správně (proč?). Nyní také jsme s počtem šestek „blíže“ k zamítnutí než při stejném poměru a méně pokusech.

Pokračujeme v pokusech dále a předpokládáme, že v dalších 120 pokusech opět napozorujeme 24 šestek. Celkem máme 480 pokusů a místo očekávaných 80 šestek máme napozorováno 96. Jak nyní dopadne hypotéza? Snadno spočítáme, že $k = 16 \doteq 1,96 \cdot 10 \cdot \sqrt{2}/\sqrt{3} = 16,003$. Hypotézu tedy zamítneme při nejvýše 74 či nejméně 96 šestkách. Což je náš případ a hypotézu o spravedlivé kostce na 5 % hladině zamítáme. Připomeňme, že podíl šestek ve výsledku pokusu je opět stejný, tedy 1/5.

V čem je předložený příklad poučný? Při zvětšujícím se počtu pozorování n se snižuje odchylka podílu p_n úspěchů (šestek) od hypotetické pravděpodobnosti (1/6), kterou ještě „tolerujeme“ ve smyslu nezamítnutí hypotézy (při pevné hladině testu). Stejný vliv na zamítnutí hypotézy má také zvýšení hladiny testu. Předvedený příklad sice nemůže nahradit rigorózní důkaz, nicméně ilustruje skutečnost, že v tomto případě (a mnoha dalších) *síla testu proti pevné alternativě roste s počtem pozorování*. Můžeme tuto diskusi uzavřít tím, že podle požadované síly testu proti dané alternativě volíme hladinu testu a rozsah výběru. Typicky nedostaneme jednoznačné „doporučení“ a proto opět volíme kombinaci hladiny testu a počtu pokusů tak, aby byla pro nás optimální. Bohužel, ne vždy máme takovou volnost, že bychom mohli provést požadovaný počet pokusů, často je celkový možný počet pokusů omezen časovými či ekonomickými možnostmi.

6. Spojitá rozdělení

Před dalším výkladem si dovolíme jednu malou odbočku. V pravděpodobnosti a statistice se často používají spojité náhodné veličiny jako vhodný model pro modelování náhodných veličin teoreticky nabývajících libovolné hodnoty v nějakém (i nekonečném) intervalu. U takových náhodných veličin si již nevystačíme s konečnou aditivitou pravděpodobnosti. Již si také nevystačíme s přidělením pravděpodobnosti jednotlivým bodům, ta je ostatně u spojitých náhodných veličin rovna nule. Namísto toho zde definujeme pro náhodnou veličinu X pravděpodobnost výskytu v intervalu $[a, b]$ pomocí

$$P[a \leq X \leq b] = \int_a^b f(x)dx, \quad (33)$$

kde f je nezáporná funkce taková, že $\int_{-\infty}^{\infty} f(x)dx = 1$. Této funkci se říká hustota rozdělení veličiny X . Pomocí hustoty můžeme počítat nejenom prav-

děpodobnost, ale i střední hodnotu

$$EX = \int_{-\infty}^{\infty} x f(x) dx \quad (34)$$

a další charakteristiky díky předpisu

$$Eg(X) = \int_{-\infty}^{\infty} g(x) f(x) dx \quad (35)$$

pro ty funkce, pro něž má integrál smysl (například spojitě funkce s konečným integrálem, nebo indikátory množin).

Významné postavení pro testování hypotéz mají *kvantily* rozdělení. Kvantil q_α definujeme pro $\alpha \in (0, 1)$ jako řešení rovnice

$$\int_{-\infty}^{q_\alpha} f(x) dx = P[X \leq q_\alpha] = \alpha. \quad (36)$$

Kvantil q_α rozděluje možné hodnoty náhodné veličiny na dvě části. Hodnot menších než q_α nabývá X s pravděpodobností α , hodnot větších než q_α s pravděpodobností $1 - \alpha$. Důležitou úlohu při popisu dat představuje medián $\tilde{X} = q_{1/2}$, střední kvantil. Kvantil obecně nemusí být dán jednoznačně, v tom případě se volí zpravidla nejmenší hodnota splňující definiční rovnost.

Nyní již můžeme přikročit k testování parametru normálního rozdělení.

7. Testy o střední hodnotě normálního rozdělení

Již jsme viděli, jak aproximace normálním rozdělením umožní snadno testovat hypotézu o pravděpodobnosti úspěchu při mnoha provedených nezávislých pokusech z alternativního rozdělení. Ukážeme si, jak testovat hypotézu o střední hodnotě normálního rozdělení. Připomeňme přitom, že náhodná veličina X má normální rozdělení o střední hodnotě μ a rozptylu σ^2 , pokud

$$\frac{X - \mu}{\sigma} \sim N(0, 1), \quad (37)$$

neboli $X = \mu + \sigma U$, kde náhodná veličina U se řídí normovaným normálním rozdělením $N(0, 1)$. Ke studiu normálního rozdělení nám tak stačí znalost normovaného normálního rozdělení z níž odvodíme vše potřebné. Jedná se tedy o třídu rozdělení uzavřenou vůči lineární transformaci, kde $\mu \in \mathbb{R}$ a $\sigma > 0$. Tyto dva parametry také normální rozdělení plně určují. Normální rozdělení je modelem užívaným pro mnoho případů spojitých náhodných veličin, uveďme například výšku či IQ u lidí.

Omezíme se zde pouze na test o střední hodnotě. Zbývá však rozptyl, který v praxi většinou neznáme a tvoří tak rušivý parametr našeho testu.

Předpokládejme nyní, že máme k dispozici náhodný výběr z normálního rozdělení $X_1, \dots, X_n$. Všechny tyto veličiny mají stejnou střední hodnotu μ , stejný rozptyl σ^2 a jsou nezávislé. Pro normální rozdělení platí, že součet normálně rozdělených veličin má opět normální rozdělení. Z toho důvodu má výběrový průměr $\bar{X}$ též normální rozdělení, tentokrát se střední hodnotou μ a rozptylem σ^2/n . Této skutečnosti využijeme při tvorbě testu o střední hodnotě.

7.1. Test při známém rozptylu

Nyní je to snadné. Předpokládejme, že máme hypotézu s alternativou

$$H_0 : \mu = \mu_0 \text{ proti } H_1 : \mu \neq \mu_0. \quad (38)$$

Proti hypotéze hovoří malé i velké hodnoty výběrového průměru. O něm víme, že za platnosti hypotézy má střední hodnotu μ_0 a

$$Z = \sqrt{n} \frac{\bar{X} - \mu_0}{\sigma} \sim N(0, 1). \quad (39)$$

Záleží tak na normované hodnotě $|Z|$, která, abychom měli důvod k zamítnutí hypotézy, musí překročit kritickou hodnotu. Platí

$$P[|Z| > k] = 2\Phi(-k), \quad k > 0. \quad (40)$$

Při dané hladině α je snadné zjistit, zda máme hypotézu zamítnout či nikoliv. Hledáme tedy k tak, aby $P[|Z| > k] = \alpha$. Například pro hladinu 5 % je $k = 1,96$, což odpovídá kvantilu $q_{0,975}$ pro normální rozdělení.

Někdy je potřeba testovat jednostrannou hypotézu proti jednostranné alternativě

$$H_0 : \mu \leq \mu_0 \text{ versus } H_1 : \mu > \mu_0. \quad (41)$$

Obdobným způsobem přijdeme k závěru, že proti hypotéze svědčí vysoké hodnoty Z (tentokrát bez absolutní hodnoty) a hledáme kritickou hodnotu tak, že $P[Z > k] = 1 - \Phi(k) = \alpha$. Zde pro hladinu 5 % dostáváme $k = 1,645$. Všimněme si, že tuto hypotézu zamítáme již pro menší kladné hodnoty Z než oboustrannou, zatímco ji nezamítáme vůbec pro záporné hodnoty Z . Proto je tento test zásadně doporučován v případech, kdy potřebujeme testovat jednostrannou hypotézu, je totiž citlivější. Je *zásadně nutné* mít rozhodnutí o použití jednostranné či oboustranné alternativy připravené před samotným

testem, *nesmíme* přizpůsobovat hypotézu a alternativu napozorovanému výběru.

Test opačné jednostranné hypotézy odvodíme snadno.

Příklad 4. Chceme zjistit, zda průměrná výška osmnáctiletých nastupujících studentů na UK je 176 cm, tedy zda odpovídá průměrné výšce v populaci. Test máme provést na hladině 5 %. Náhodně vybraných 168 nastupujících studentů bylo změřeno a jejich průměrná výška je 175,2 cm. Populační rozptyl je známý a je $\sigma^2 = 67,4$. Testujeme tedy proti oboustranné alternativě zahrnující obě odchýlení od průměru. Spočítejme

$$|Z| = 12,96 \frac{|175,2 - 176|}{8,21} = 1,2629 \quad (42)$$

Protože tato hodnota je menší než 1,96, hypotézu o průměrné výšce 176 cm nezamítáme.

7.2. Test při neznámém rozptylu

Odvození testu je obdobně snadné pomůžeme-li si zde nedokazovaným výsledkem z literatury.

Lemma 2 (t-rozdělení). *Pro náhodný výběr o rozsahu $n \geq 2$ z normálního rozdělení se střední hodnotou μ má náhodná veličina*

$$T = \sqrt{n} \frac{\bar{X} - \mu}{S}, \quad (43)$$

kde S^2 je nestranný výběrový rozptyl, (Studentovo) t rozdělení o $n-1$ stupních volnosti.

Všimněme si chybějícího parametru σ , který byl nahrazen v 39 odhadem S . Nicméně zdůvodnění předchozího výsledku je poněkud hlubší a zájemce odkazujeme na uvedenou studijní literaturu.

Nyní již vše proběhne stejně jako u testu se známým rozptylem, jen s tím rozdílem, že

$$P[|T| > k] = 2T_{n-1}(-k), \quad k > 0 \quad (44)$$

kde T_{n-1} je distribuční funkce Studentova rozdělení o $n-1$ stupni volnosti. Uveďme si například, že pro test na hladině 5 % a rozsah výběru 16 je kritická hodnota, tedy kvantil $q_{0,975}$, rovna $k = 2,131$. Obecně jsou tyto kritické hodnoty vyšší než odpovídající kritické hodnoty normálního rozdělení, s rostoucím n se k nim však monotónně blíží.

Příklad 4 (pokračování). Necht populační rozptyl výšky osmnáctiletých mužů není znám a my jsme odkázáni na jeho odhad $S^2 = 56,85$ naměřený na výběru 168 studentů. Pak

$$|T| = 12,96 \frac{|175,2 - 176|}{7,54} = 1,375. \quad (45)$$

Kritickou hodnotou na hladině 5 % pro tento test je $k = 1,975$ jíž naše statistika $|T|$ nepřekračuje. Na hladině 5 % proto nezamítáme hypotézu o tom, že by průměrná výška nastupujících studentů UK byla 176 cm.

8. Problém dvou výběrů

Testování hypotézy o parametru rozdělení náhodné veličiny je zajímavý problém, nicméně ještě zajímavější je otázka porovnání dvou veličin. Omezme se na ukázkou testu pro spojitě náhodné veličiny.

Jen pro úplnost uvedme, že základním předpokladem porovnání dvou náhodných výběrů (tedy hypotéza o dvou populacích) je nezávislost výběrů. Dalším, ne vždy nezbytným, předpokladem je stejný model pro sledované náhodné veličiny. Tento model se může lišit parametrem, který je stejný v rámci jedné populace, ale mezi sebou se populace lišit mohou. Populace se mohou lišit například svou střední hodnotou či rozptylem, což jsou nejčastější uvažované rozdíly. Mějme vždy na mysli, že bychom měli porovnávat porovnatelné, a že výsledky by měly mít rozumný smysl v oblasti, v níž statistiku aplikujeme. Také interpretace výsledků je důležitou součástí analýzy a ukvapené závěry mohou, byť „podložené tvrdými statistickými údaji“, být naprosto zcestné. Zde je na místě opět zdůraznit význam pečlivé přípravy experimentu. Vymýšlet hypotézy, otázky a odpovědi nad naměřenými daty již bývá pozdě. Mít připravené otázky i požadavky na to, co od statistiky vlastně v daném problému očekáváme, připravené odpovědi (kladné či záporné) a interpretace, rozmyšlenou hypotézu, hladinu a požadovanou sílu, to vše mnohokrát zvyšuje vypovídající hodnotu výzkumu.

Pro porovnání dvou spojitých náhodných veličin je možné využít uspořádání všech výsledků pokusu podle velikosti a porovnat pořadí výsledků v obou zkoumaných výběrech. Takovým testům založeným na pořadí se říká neparametrické, protože jsou nezávislé na konkrétním rozdělení spojitě náhodné veličiny. Ukážeme si provedení *Wilcoxonova testu*. Předpokladem Wilcoxonova testu je, že rozdělení (neznámé, spojitě) obou náhodných výběrů je stejné až na posunutí. Obecně se tedy dá říci, že náhodná veličina X odpovídající prvnímu výběru má stejné rozdělení jako náhodná veličina $Y + \theta$, kde Y odpovídá druhému výběru a θ je neznámá hodnota.

Příklad 6. Zkoumáme, zda průměr známek na vyvědčení se liší u chlapců a děvčat. K tomu účelu vybereme v jednom populačním ročníku osm chlapců a pět děvčat a porovnáme jejich průměry na vysvědčení. Hypotéza je, že průměrné výsledky se u chlapců a děvčat neliší, testovat máme na hladině 0,05.

Jak budeme postupovat? Různé počty vzorků nejsou překážkou pro test hypotézy, pokud nejde o řádové rozdíly. Představme si, že výsledky prvního výběru označíme X_i , $i = 1, \dots, m$. Výsledky druhého výběru označíme Y_i pro $i = 1, \dots, n$. Celkem máme $m + n$ pozorování, která jsou s pravděpodobností 1 různá za předpokladu spojitosti sledované náhodné veličiny.

Nyní seřadíme všechna pozorování X_i a Y_i dle velikosti a přiřadíme jim jejich pořadí v uspořádaném a setříděném souboru R_i , $i = 1, \dots, m + n$, číslo mezi 1 a $m + n$. Pro jednoduchost předpokládejme, že naměřené hodnoty jsou vesměs různé. Pokud by tomu tak nebylo, dáváme shodným hodnotám jejich průměrné pořadí. Nejmenší pozorování bude mít hodnotu 1, největší $m + n$. Dostáváme tak

$$\begin{array}{cccccccc} X_1 & X_2 & \dots & X_m & Y_1 & Y_2 & \dots & Y_n \\ R_1 & R_2 & \dots & R_m & R_{m+1} & R_{m+2} & \dots & R_{m+n} \end{array} \quad (46)$$

Celkem je součet všech pořadí roven $(m + n)(m + n + 1)/2$. Testujeme-li hypotézu o shodě rozdělení, tedy $\theta = 0$, očekáváme, že

$$Z_X = \sum_{i=1}^m R_i \doteq \frac{m(m + n + 1)}{2}, \quad Z_Y = \sum_{i=m+1}^{m+n} R_i \doteq \frac{n(m + n + 1)}{2}, \quad (47)$$

tedy součet pořadí odpovídajících prvnímu a druhému výběru bude přibližně ve „správném“ poměru. Pokud je tento součet pořadí příliš vzdálený od očekávané hodnoty, jsme rozhodnutí hypotézu zamítnout.

Příklad 6 (pokračování). Podívejme se na seřazené výsledky žáků i s celkovým pořadím.

Chlapci	1,08	1,58	1,67	1,83	1,92	2,00	2,17	2,50
	2	6	7	9	10	11	12	13
Đevčata	1,00	1,17	1,33	1,50	1,75			
	1	3	4	5	8			

Sečteme-li pořadí odpovídající děvčatům dostaneme hodnotu 21 ($= 1 + 3 + 4 + 5 + 8$). Celkový součet pořadí je 91 ($= 13 \cdot 14/2$). Přitom při platnosti nulové hypotézy bychom očekávali součet pořadí pro děvčata asi 35 ($= 5 \cdot 14/2$).

Hodnota 21 je dle tabulek pro Wilcoxonův dvouvýběrový test pro hladinu 0,05 již významná. Hypotézu o shodných školních výsledcích děvčat a chlapců tedy na této hladině zamítneme.

U průměru na vysvědčení normalitu předpokládat nemůžeme, neboť zde nemáme jednoznačnou oporu ve zkušenosti. Na druhou stranu u IQ normalitu předpokládat můžeme, neboť IQ je zavedeno tak, aby v populaci mělo rozdělení co nejvíce shodné s normálním. Pak pro porovnání IQ dvou populací lze použít postupu podobného jako pro test o neznámém parametru střední hodnoty normálního rozdělení. Jedná se pak o alternativní test k Wilcoxonovu testu, takzvaný *dvouvýběrový t-test*. Také dvouvýběrový t-test je testem proti alternativě posunutí.

Testů porovnávajících dva výběry je celá řada, tak jako je řada rozdílů které chceme případně odhalit. Podle toho pak volíme takový test, který je k danému rozdílu nejcitlivější. Pro stejný problém můžeme obvykle vybírat z více testů. Pak většinou záleží na modelu, který test je vhodnější. Například při alternativě posunutí je lepší použít dvouvýběrový t-test, pokud máme data normálně rozdělená. Pokud máme o normalitě pochyby, nebo víme že data být normálně rozdělena nemohou, je lepší dvouvýběrový Wilcoxonův test. Pro diskrétní pozorování pak bývá doporučován dvouvýběrový znaménkový test, což je obměna Wilcoxonova testu.

8.1. Statistická významnost

Poslední, co nám v této části zbývá diskutovat je otázka: „Jak poznáme *statisticky významnou* odchylku“. Pro daný rozsah výběru jde o funkci hladiny testu. Pro menší rozsahy výběru se dají najít tabulky přesných kritických hodnot, dnes jsou běžnou součástí statistických programů. Pro větší rozsahy pak většinou spoléháme na nějaké asymptotické chování testovací statistiky, často zejména asymptotickou normalitu (nezapomeňme však na výjimky, jako je například aproximace binomického rozdělení Poissonovým pro velké rozsahy výběru a malé hodnoty p). Pro asymptotické testy používáme kritické hodnoty pro normální rozdělení, které jsou běžnou součástí statistických programů a tabulek. Připomeňme však, že dnes již málokdy provádíme výpočty testů ručně. Naopak obvykle použijeme vhodný program a ten nám, po vybrání požadovaného testu, dodá veškeré nezbytné výsledky. Včetně toho, že pokud to náš program a počítač zvládnou, dostaneme výsledky přesného testu, jinak dostaneme nabídnuté asymptotické výsledky.

Statistická významnost je ve speciálních případech pěkně vidět na intervalu spolehlivosti.

9. Interval spolehlivosti

Obecně uvažovaný problém spolehlivostní množiny pro parametr rozdělení vytváří velmi složitou matematickou úlohu. Jde vlastně o úlohu najít na základě náhodného výběru $\mathbf{X}$ podmnožinu $\Theta(\mathbf{X})$ parametrického prostoru tak, aby tato množina obsahovala neznámý parametr s danou pravděpodobností a zároveň byla nějak optimální. Zde si však uvedeme jen speciální příklad pro normální rozdělení a jeho vztah k testování hypotéz.

Lze říci, že pro úlohy uvedené výše můžeme stanovit spolehlivostní množinu (zde půjde o intervaly) takto. *Interval spolehlivosti* pro parametr θ na hladině *spolehlivosti* $1 - \alpha$ je množina $\Theta_0 \subset \Theta$ těch parametrů θ' , že test hypotézy $\theta = \theta'$ proti oboustranné alternativě nezamítneme na hladině α . Tento interval je zřejmě dán realizací $\mathbf{x}$ náhodného výběru.

Již z uvedené definice je zřejmé, že mezi intervalem spolehlivosti a testem hypotézy proti oboustranné alternativě je jednoznačný vztah. Obdobně bychom mohli zavést jednostranné intervaly spolehlivosti odpovídající testům hypotézy proti jednostranné alternativě.

Pojďme se teď věnovat speciálnímu případu normálního rozdělení. Buď $\mathbf{X} = (X_1, \dots, X_n)$ náhodný výběr o rozsahu n z normálního rozdělení o neznámém rozptylu. Již víme, že test hypotézy pro střední hodnotu $H_0 : \mu = \mu_0$ proti oboustranné alternativě $H_1 : \mu \neq \mu_0$ na hladině α je založen na hodnotě statistiky

$$Z = \sqrt{n} \frac{\bar{X} - \mu_0}{S}, \quad S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2.$$

Na hladině α hypotézu zamítneme, pokud $|Z| > u_{1-\alpha/2}$, kde u_α značí α kvantil normovaného normálního rozdělení.

Vyjděme z definice, že interval spolehlivosti na hladině (spolehlivosti) $1 - \alpha$ je množina všech μ takových, že nezamítneme hypotézu H_0 na hladině (testu) α . Pak tedy hledáme takové hodnoty μ , že

$$\begin{aligned} |Z| \leq u_{1-\alpha/2} &\Leftrightarrow \sqrt{n} \frac{\bar{X} - \mu}{S} \leq u_{1-\alpha/2} \\ &\Leftrightarrow \mu \in [\bar{X} - u_{1-\alpha/2} S / \sqrt{n}, \bar{X} + u_{1-\alpha/2} S / \sqrt{n}]. \end{aligned} \quad (48)$$

Tím máme interval spolehlivosti stanoven. Všimněme si jeho závislosti na výběru prostřednictvím $\bar{X}$ a S , jeho závislosti na hladině spolehlivosti prostřednictvím $u_{1-\alpha/2}$ a jeho závislosti na rozsahu výběru prostřednictvím $\sqrt{n}$.

Zejména poznamenejme, že šířka intervalu spolehlivosti je *přímo úměrná* výběrovému rozptylu S^2 , *přímo úměrná* hladině spolehlivosti $1 - \alpha$ (a tak *nepřímo úměrná* hladině testu α) a *nepřímo úměrná* rozsahu výběru n .

Příklad 4 (pokračování). Chceme stanovit interval spolehlivosti pro průměrnou výšku nastupujících studentů UK na hladině spolehlivosti 0,95. Máme $\bar{X} = 175,2$, $S = 8,21$ a $\sqrt{n} = 12,96$. Použitím (48) dostaneme interval spolehlivosti

$$[175,2 - 8,21 \cdot 1,975/12,96, 175,2 + 8,21 \cdot 1,975/12,96] = [173,95, 176,45]. \quad (49)$$

Všimněme si, že testovaná hodnota 176 leží ve zjištěném intervalu spolehlivosti a proto jsme nemohli zamítnout $H_0 : \mu = 176$ na hladině 5 %. Vidíme také, které všechny hypotézy bychom nezamítli. Zajímavá je také šířka intervalu spolehlivosti na základě 168 pozorování. Jde o 2,5 centimetru.

Zde je na místě poznámka o interpretaci intervalu spolehlivosti. Jelikož interval spolehlivosti má *náhodné meze*, zatímco neznámý parametr je *pevnou hodnotou*, neleží parametr v intervalu spolehlivosti, ale naopak. *Interval spolehlivosti obsahuje neznámý parametr s pravděpodobností $1 - \alpha$* . V tomto světle je interval spolehlivosti *intervalovým odhadem* parametru, na rozdíl od bodového odhadu uvedeného v části 2.

10. Závěrem

Statistické uvažování je založeno na znalosti matematických vlastností předpokládaného pravděpodobnostního modelu. Na základě zkušeností či teorie předpokládáme pro náš experiment nějaký model „řídící náhodu“. Ukázali jsme si, že náhodný výběr je modelem, v jehož rámci se dá provést bohatá statistická analýza. Náhodný výběr skutečně je modelem, neboť předpokládáme nezávislost a stejné rozdělení. Proto se snažíme o to, aby náš experiment nezávislost a stejné rozdělení pokusů zaručoval.

Jak jsme viděli, speciálním případem pravděpodobnostních modelů pro náhodný výběr jsou takzvané parametrické modely. Typ (třída, rodina) rozdělení je známý, specifikovaný nějakým matematickým předpisem obsahujícím hodnotu náhodné veličiny a parametr. Hodnota parametru je však neznámá. Úkolem statistiky je pak hledat odhady tohoto parametru a rozhodovat o hypotézách o tomto parametru.

Předvedli jsme si však i příklad testu hypotézy, kde třída rozdělení nebyla specifikována, jde o Wilcoxonův test. Místo předpokladu o tvaru rozdělení jsme se omezili na předpoklad, že jde o rozdělení spojitě náhodné veličiny, čímž jsme vylučovali (alespoň teoreticky) výskyt shodných pozorování. Předpokládali jsme ovšem, že oba porovnávané výběry mají shodný tvar rozdělení lišící se nejvýše v posunutí o konstantu. Takovým postupům říkáme neparametrické.

Zdůrazněme, že volba modelu, zde nediskutovaná, má na kvalitu statistických závěrů významný vliv. Zvolíme-li například model normálního rozdělení pro malý výběr značně nesymetrických dat (čili se nemůžeme odvolávat ani na asymptotické chování průměru), pak závěr byť sebezpečněji spočítaného testu na pečlivě zvolené hladině bude stejně k ničemu. Automatické používání normálního rozdělení na spojitá data (s poměrně vágním zdůvodněním, že se „to tak dělá“) nedoporučujeme. V případě nejasností je nakonec lepší použít neparametrickou metodu, než sice optimální parametrickou metodu, ovšem na data nesplňující její předpoklady.

10.1. Statistická a věcná významnost

Další, často opomíjená, záludnost statistického rozhodování spočívá v rozdílu mezi statistickou a věcnou významností. Nechtě hypotetická testovaná hodnota parametru je rovna θ_0 a my chceme odhalit neplatnost této hypotézy s vysokou pravděpodobností, pokud se pro skutečný parametr θ platí $|\theta - \theta_0| > \eta$. Hodnota η zde reprezentuje *věcně významný rozdíl*, tedy rozdíl, který považujeme za významný ne statisticky, ale z hlediska jeho důsledků (významný pro život, experiment, ...). Stává se, že pro malé rozsahy výběrů statistický test není (a ani nemůže být) takového odhalení schopen a síla takového testu je malá i pro rozdíl $|\theta - \theta_0| > \eta$, ba dokonce někdy i mnohem větší než η . Na druhou stranu, pro velké rozsahy výběru se zákonitě stává, že pro fakticky zanedbatelný rozdíl $|\theta - \theta_0|$ máme test s vysokou silou a hypotézu zamítáme, ačkoliv po věcné stránce nás takový rozdíl vlastně nezajímá. Představme si *statisticky prokázaný* rozdíl mezi obvodem hlavy novorozenců v Praze a v Brně 0,03 cm. Statisticky je takový rozdíl možná významný, ale *věcně zcela nevýznamný*, tedy v reálném životě či medicíně jde o nezájímavý výsledek.

Pěkně je vidět vztah statistické a věcné významnosti u intervalu spolehlivosti. Dokonce bychom mohli říci, že v případech, kdy interval spolehlivosti je ekvivalentní testu hypotézy (jako je tomu u normálního rozdělení a to i v případě asymptotických testů), doporučujeme „testovat“ hypotézu na základě zjištění, zda je testovaná hodnota obsažena v intervalu spolehlivosti. V případě dvouvýběrového testu (a odpovídajícího intervalu spolehlivosti – viz například uvedená literatura) jsou všechny rozdíly mimo interval spolehlivosti statisticky významné. Můžeme si také sami vytvořit interval „věcné významnosti“ a porovnat jej pro názornost s intervalem statistické významnosti.

Nejsme-li spokojeni se silou testu proti věcně významným rozdílům, potřebujeme zvýšit počet pozorování. Druhou možností, tedy zvýšení pravděpodobnosti chyby prvního druhu nemůžeme příliš využívat. Chceme ji totiž

držet pod kontrolou a nízkou. Jsme-li naopak v situaci, kdy by i věcně nevýznamný rozdíl vyšel statisticky významný, můžeme snížit hladinu testu tak, aby test měl velkou sílu až pro opravdu věcně významné rozdíly. Díky velkému počtu pozorování bychom tak mohli snížit pravděpodobnost chyby prvního druhu na zanedbatelnou hodnotu a přitom mít pro věcně významný rozdíl dostatečně vysokou sílu. Bohužel, běžně se takový postup ke škodě uživatelů i statistiky nepoužívá

Literatura

- [1] Anděl J. (1985). *Matematická statistika*, SNTL/Alfa, Praha.
- [2] Anděl J. (2005). *Základy matematické statistiky*, MATFYZPRESS, Praha.
- [3] Zvára K. a Štěpán J. (2003). *Pravděpodobnost a matematická statistika*, MATFYZPRESS, Praha.

REGRESE

Karel Zvára

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

karel.zvara@mff.cuni.cz

1. Úvod

Regresní model patří k nejčastěji používaným statistickým modelům. Vyjadřujeme jím učeně řečeno podmíněnou střední hodnotu *vysvětlované (závislé) proměnné* jako funkci jiné či jiných proměnných, nazývaných podle okolností jako regresory či *nezávislé proměnné*. *Regresní funkci* se pak rozumí $\mu(x) = E(Y|X = x)$. Regresní funkce tedy udává, jaká je střední hodnota náhodné veličiny Y při dané hodnotě x náhodného vektoru $\mathbf{X}$. Vektor x můžeme zvolit sami, a to pevně, nenáhodně. Pak chápeme $\mathbf{X}$ jako speciální náhodný vektor s nulovým rozptylem.

Nejznámějším speciálním případem je funkce $\mu(x) = \beta_0 + \beta_1 x$. Závislost je pak popsána *regresní přímkou* $y = \beta_0 + \beta_1 x$.

Samotné označení *regrese* pochází z předminulého století, kdy se sir Francis Galton zabýval závislostí výšky syna na výšce otce. Vezměme otce, kteří se liší svojí výškou řekněme o 10 cm od průměrné výšky otců. Galton zjistil mimo jiné, že se synové takových otců se od průměrné výšky synů liší v průměru jen asi o 5 cm. Je patrný návrat zpět k průměru, tedy regrese.

Kolega Saxl byl tak laskav a upozornil mne s odkazem na [2], že se úlohou dědičnosti výšky mužů zabýval již kolem roku 1730 Sir John Gordon of Park. Tehdy však ještě termín *regrese* nepoužil.

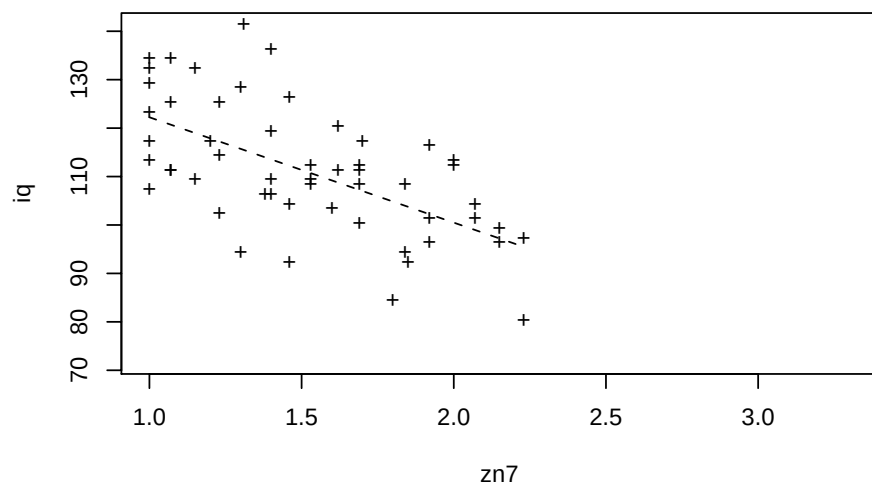
Možnosti regresního modelování si ukážeme na datech, která by mohla pocházet z nějaké diplomové práce, která se zabývá souvislostí školních výsledků s hodnotou inteligenčního koeficientu IQ.

2. Jak souvisí IQ se školním prospěchem?

Vezměme údaje o 54 dívkách ze sedmé třídy. Obrázek 1 naznačuje, že v průměru s rostoucím průměrem známek zpravidla klesá hodnota IQ. Lze tuto závislost nějak popsat a prokázat?

Klíčová slova: Náhodný výběr, bodový odhad, test statistické hypotézy.

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839.



Obrázek 1. Souvislost IQ s průměrným prospěchem v 7. třídě u dívek.

Body na obrázku je vedena přímka, která data v jakémsi přesně definovaném smyslu vystihuje nejlépe. Z příslušného programu vypadne vztah $y = 143,987 - 21,751 \cdot x$, kde y je běžně užívané označení pro vysvětlovanou veličinu (zde IQ) a x označení pro nezávisle proměnnou (zde pro průměrnou známku).

Uvedený vztah má jednoduchou interpretaci. Máme-li odhadnout IQ žákyně, která má průměr 1, dostaneme po dosazení odhad $143,987 - 21,751 \cdot 1 = 122,236$. Pro dvojkařku ovšem dostaneme odhad jen $143,987 - 21,751 \cdot 2 = 100,485$.

3. Metoda odhadu

Nyní vysvětlíme, podle jakého principu jsme vlastně určili odhad koeficientů přímky. Výchozí představou je *model* uvažované závislosti, totiž vztah ($1 \leq i \leq n$)

$$Y_i = \beta_0 + \beta_1 x_i + E_i, \quad (1)$$

kde Y_i jsou realizace náhodné veličiny IQ, x_i jsou známé konstanty (školní průměry) a platí $n > 2$. Neznámé parametry β_0, β_1 jsou součástí systematické složky IQ, náhodné veličiny $E_1, \dots, E_n$ popisují náhodnou složku IQ.

Předpokládá se, že to jsou *nezávislé* náhodné veličiny s nulovou střední hodnotou a rozptylem (pro všechny stejným) σ^2 . Zpravidla se předpokládá také normální rozdělení. V příkladu s měřením IQ lze tento předpoklad považovat za splněný. Ukazatel IQ byl konstruován s cílem dosáhnout normálního rozdělení.

Modelové vztahy (1) můžeme pro všechna n současně zapsat jako

$$\begin{pmatrix} Y_1 \\ Y_2 \\ \vdots \\ Y_n \end{pmatrix} = \beta_0 \begin{pmatrix} 1 \\ 1 \\ \vdots \\ 1 \end{pmatrix} + \beta_1 \begin{pmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{pmatrix} + \begin{pmatrix} E_1 \\ E_2 \\ \vdots \\ E_n \end{pmatrix}. \quad (2)$$

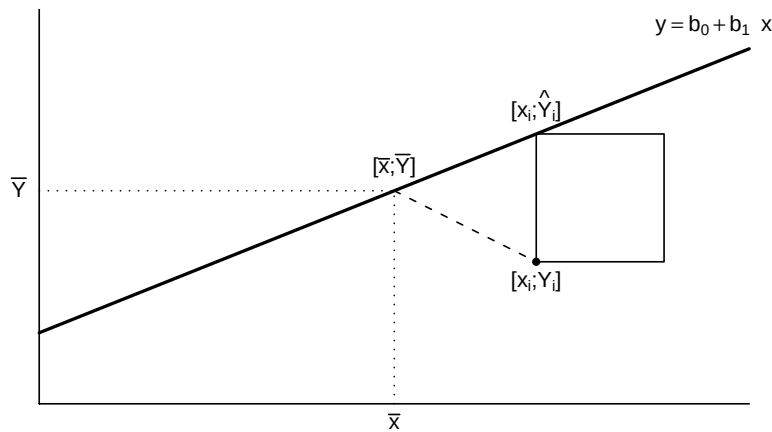
Pozorované hodnoty vysvětlované proměnné (hodnoty IQ) si představme jako náhodný vektor $\mathbf{Y}$ o složkách $Y_1, \dots, Y_n$, vektor hodnot nezávisle proměnné (průměrných známek) označíme symbolem $\mathbf{x}$, vektor n jedniček jako $\mathbf{1}$. Vztah (2) pak můžeme vektorově zapsat jako $\mathbf{Y} = \beta_0 \mathbf{1} + \beta_1 \mathbf{x} + \mathbf{E}$. Vektor $\mathbf{E}$ je náhodná složka vyjadřující individuální nahodilé odchylky skutečné hodnoty IQ od složky systematické. Náhodná veličina E_i vyjadřuje odchylku skutečné hodnoty IQ i -té dívky od střední hodnoty IQ všech dívek se stejným školním prospěchem. Velikostí písmen označujících vyšetřované proměnné rozlišujeme jejich kvalitu. Pro dané (tedy už nenáhodné) hodnoty x_i (školní prospěch) se snažíme odhadnout střední hodnoty náhodných veličin Y_i (hodnot IQ).

Výpočet odhadů parametrů β_0, β_1 si můžeme představit tak, že vlastně hledáme takovou lineární kombinaci vektorů $\mathbf{1}$ a $\mathbf{x}$, která je k realizaci náhodného vektoru $\mathbf{Y}$ nejbližší. Náhodnou složku $\mathbf{E}$ vyjadřující nahodilé odchylky od pravidla totiž potřebujeme vyloučit, protože má nulovou střední hodnotu. Odhady parametrů β_0, β_1 , které označíme b_0, b_1 , dostaneme minimalizací funkce (čtverec vzdálenosti $\mathbf{Y}$ od lineární kombinace vektorů $\mathbf{1}$ a $\mathbf{x}$)

$$S(\beta_0, \beta_1) = \|\mathbf{Y} - (\beta_0 \mathbf{1} + \beta_1 \mathbf{x})\|^2 = \sum_{i=1}^n (Y_i - \beta_0 - \beta_1 x_i)^2.$$

Odtud název metody odhadu – *metoda nejmenších čtverců*. Odhad metodou nejmenších čtverců má mnohé optimální vlastnosti, zejména tehdy, když mají náhodné veličiny E_i normální rozdělení.

Na schematickém obrázku obr. 2 je znázorněn jeden takový čtverec. Nepřehlédněte, že strany čtverce jsou rovnoběžné s osami souřadnic. Strana čtverce je dána vzdáleností bodu $[x_i, Y_i]$ se skutečnou hodnotou náhodné veličiny Y_i od jeho protějšku $[x_i, \hat{Y}_i]$ na ležícího na přímce, který má **stejnou**



Obrázek 2. Znázornění metody nejmenších čtverců. Čárkovaná úsečka spojuje body $[x_i, Y_i]$ a $[\bar{x}, \bar{Y}]$, u bodu $[x_i, Y_i]$ je nakreslen čtverec s plochou $(Y_i - b_0 - b_1 \cdot x_i)^2$.

hodnotu x -ové souřadnice, tj. $\hat{Y}_i = b_0 + b_1 x_i$. Nejde tedy o vzdálenost bodu od přímky známou z geometrie. Přímka $y = b_0 + b_1 x$ má vlastnost, že celková plocha n takovýchto čtverců je nejmenší.

Zmíněná minimalizace je pěkným cvičením na extrém funkce dvou proměnných. Parciální derivace podle β_0 a β_1 jsou

$$\begin{aligned}\frac{\partial}{\partial \beta_0} \mathcal{S}(\beta_0, \beta_1) &= -2 \sum_{i=1}^n (Y_i - \beta_0 - \beta_1 x_i), \\ \frac{\partial}{\partial \beta_1} \mathcal{S}(\beta_0, \beta_1) &= -2 \sum_{i=1}^n (Y_i - \beta_0 - \beta_1 x_i) x_i.\end{aligned}$$

Když tyto derivace anulujeme a předpokládáme $\sum_{i=1}^n (x_i - \bar{x})^2$ (tj. mezi hodnotami x_i jsou aspoň dvě různé), dostaneme řešením vzniklé soustavy dvou lineárních rovnic odhady parametrů β_1 a β_0 ve tvaru

$$b_1 = \frac{\sum (x_i - \bar{x})(Y_i - \bar{Y})}{\sum (x_i - \bar{x})^2}, \quad (3)$$

$$b_0 = \bar{Y} - b_1 \bar{x}. \quad (4)$$

Že takto vyjde právě minimum plyne z toho, že jde o minimum nezáporné kvadratické funkce v β_0, β_1 .

Zamysleme se nad odhadem b_1 . Jedná se o lineární kombinaci nezávislých náhodných veličin Y_i s normálním rozdělením, takže také b_1 je náhodná veličina s normálním rozdělením. Pokusme se dále zjistit, jaký vztah má tento odhad k odhadovanému parametru β_1 . Podle (1) víme, že pro střední hodnotu vysvětlované proměnné Y_i platí

$$EY_i = \beta_0 + \beta_1 x_i,$$

a odtud také

$$E\bar{Y} = \beta_0 + \beta_1 \bar{x}.$$

Proto je

$$\begin{aligned} Eb_1 &= E \frac{\sum (x_i - \bar{x})(Y_i - \bar{Y})}{\sum (x_i - \bar{x})^2} \\ &= \sum_{i=1}^n \frac{(x_i - \bar{x})}{\sum_{t=1}^n (x_t - \bar{x})^2} E(Y_i - \bar{Y}) \\ &= \sum_{i=1}^n \frac{(x_i - \bar{x})}{\sum_{t=1}^n (x_t - \bar{x})^2} \beta_1 (x_i - \bar{x}) \\ &= \beta_1 \sum_{i=1}^n \frac{(x_i - \bar{x})^2}{\sum_{t=1}^n (x_t - \bar{x})^2} = \beta_1, \end{aligned}$$

takže jde o *nestranný* odhad neznámého parametru β_1 . Podobně lze spočítat, že rozptyl odhadu b_1 je

$$\text{var} b_1 = \frac{\sigma^2}{\sum_{t=1}^n (x_t - \bar{x})^2}. \quad (5)$$

Na okraj si všimněme zajímavé interpretace odhadu b_1 směrnice regresní přímky, který můžeme upravit na tvar

$$b_1 = \sum_{i: x_i \neq \bar{x}} \frac{(x_i - \bar{x})^2}{\sum_{t=1}^n (x_t - \bar{x})^2} \cdot \frac{Y_i - \bar{Y}}{x_i - \bar{x}}. \quad (6)$$

Případné vynechání sčítanců s $x_i = \bar{x}$ nic neovlivní, protože jsou stejně nulové. Vzorec (6) ukazuje, že b_1 je konvexní kombinací (váženým průměrem) směrnic úseček spojujících body $[x_i, Y_i]$ s bodem $[\bar{x}, \bar{Y}]$.

Nejmenší hodnota funkce $\mathcal{S}$, tedy $\mathcal{S}(b_0, b_1)$, se nazývá *reziduální součet čtverců*. U regresní přímky lze reziduální součet čtverců vyjádřit ve tvaru

$$\mathcal{S}(b_0, b_1) = \sum_{i=1}^n (Y_i - \bar{Y})^2 - b_1^2 \sum_{i=1}^n (x_i - \bar{x})^2, \quad (7)$$

což do jisté míry ukazuje, o kolik se výchozí variabilita vysvětlované proměnné popsaná výrazem $\sum_{i=1}^n (Y_i - \bar{Y})^2$ sníží v důsledku jejího částečného vysvětlení pomocí závislosti na hodnotách $x_1, \dots, x_n$.

Velmi často se k hodnocení toho, nakolik předpokládaná závislost vysvětluje variabilitu závisle proměnné, používá *koefficient determinace* R^2 definovaný jako

$$R^2 = 1 - \frac{\mathcal{S}(b_0, b_1)}{\sum_{t=1}^n (Y_t - \bar{Y})^2}.$$

Zjišťujeme tedy, jaký díl původní variability vysvětlované proměnné se podařilo vysvětlit. Proto se koefficient determinace často uvádí v procentech.

Ještě zbývá odhadnout parametr σ^2 . Lze dokázat, že statistika

$$S^2 = \frac{\mathcal{S}(b_0, b_1)}{n - 2} \quad (8)$$

je nestranným odhadem rozptylu σ^2 a že je (stochasticky) nezávislá na vektoru $(b_0, b_1)^T$.

4. Průkaznost závislosti

V praxi nás zajímá nejen bodový odhad neznámých parametrů, ale také jejich intervalový odhad a zejména test hypotézy, že $\beta_1 = 0$. To by totiž znamenalo, že střední hodnota EY_i nezávisí na konstantě x_i , v našem příkladu pak, že hodnota IQ nesouvisí se školním prospěchem. O této hypotéze se rozhoduje pomocí tzv. *t*-statistiky

$$T = \frac{b_1}{S} \sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}. \quad (9)$$

Hypotézu $H_0 : \beta_1 = 0$ budeme na hladině α zamítat, když bude $|T| \geq t_{n-2}(\alpha)$, kde $t_{n-2}(\alpha)$ je tzv. *kritická hodnota t-rozdělení* s $n-2$ stupni volnosti

určená tak, že ji náhodná veličina s tímto rozdělení v absolutní hodnotě překročí s pravděpodobností právě α . Kritické hodnoty bývají uváděny v učebnicích (také [3]), existují speciální knihy s tabulkami (např. [1]), počítají je statistické programy.

Pokusme se tento postup (tento tvar kritického oboru) vysvětlit. Je rozumné zamítnout hypotézu, podle které je $\beta_1 = 0$, když vyjde odhad b_1 parametru β_1 příliš daleko od nuly. Při tom je třeba vzít v úvahu velikost kolísání odhadu b_1 . Při velkém kolísání (velkém rozptylu) musíme tolerovat hodnoty odhadu velmi vzdálené od nuly.

Kdybychom znali hodnotu parametru σ , pak vzhledem k (5) by měla náhodná veličina

$$\frac{b_1}{\sqrt{\text{var} b_1}} = \frac{b_1}{\sigma} \sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}$$

normální rozdělení s jednotkovým rozptylem, za platnosti testované hypotézy dokonce s nulovou střední hodnotou. Statistika T z (9) se od tohoto vzorce liší jen tím, že neznámý parametr σ je nahrazen jeho odhadem $S = \sqrt{S^2}$. Proto má za platnosti testované hypotézy statistika T tzv. Studentovo t -rozdělení s $n - 2$ stupni volnosti (připomeňte, že jsme reziduální součet čtverců dělili právě výrazem $n - 2$, abychom dostali odhad S^2).

Ještě jeden pohled na t -statistiku z (9) je užitečný. Nahraďme ve výrazu (5) pro rozptyl odhadu b_1 neznámý parametr σ^2 jeho odhadem S^2 . Odmocnina z tohoto odhadu rozptylu b_1 se nazývá *střední chyba*. Pro tuto odmocninu se často používá označení S.E. s odkazem na odhad, o jehož variabilitě vypovídá. Označení pochází z anglického *Standard Error*. Je tedy v našem případě

$$\text{S.E.}(b_1) = \frac{S}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2}}. \quad (10)$$

Testovou statistiku z (9) lze pak psát ve tvaru

$$T = \frac{b_1}{\text{S.E.}(b_1)}. \quad (11)$$

Analogicky je možno vyjádřit testové statistiky mnohých statistických testů, např. dvouvýběrového i párového t -testu.

Pokud testovanou (nulovou) hypotézu zamítneme, říkáme, že odhad b_1 je na hladině α významný nebo že jsme na hladině α prokázali závislost. Vedle bodového odhadu ($b_1 = -21,751$) můžeme spočítat také odhad intervalový. 95% interval spolehlivosti pro β_1 obsahuje takové hodnoty β_1^0 , pro které bychom hypotézu $H_0 : \beta_1 = \beta_1^0$ na 5% hladině **nezamítli**. Budou to

tedy všechna řešení nerovnosti (srovnej s t -statistikou z (11) pro test nulové hypotézy $\beta_1 = 0$)

$$\left| \frac{b_1 - \beta_1^0}{\text{S.E.}(b_1)} \right| < t_{n-2}(\alpha).$$

Řešením je interval

$$(b_1 - \text{S.E.}(b_1) \cdot t_{n-2}(\alpha); b_1 + \text{S.E.}(b_1) \cdot t_{n-2}(\alpha)). \quad (12)$$

Odtud je patrné slovní označení pro střední chybu S.E.. Tato statistika podstatnou měrou rozhoduje o šířce intervalu spolehlivosti pro β_1 , tedy o tom, jak přesným odhadem parametru β_1 odhad b_1 ve skutečnosti je.

5. Je závislost IQ na školním prospěchu průkazná?

Při troše námahy můžeme třeba na kapesní kalkulačce spočítat, že je (se zaokrouhlením)

$$\begin{aligned} \bar{x} &= 1,511; & \bar{Y} &= 111,111; \\ \sum_{i=1}^{54} (x_i - \bar{x})^2 &= 7,471; & \sum_{i=1}^{54} (Y_i - \bar{Y})^2 &= 9\,585,333; \\ \sum_{i=1}^{54} (x_i - \bar{x})(Y_i - \bar{Y}) &= -162,499. \end{aligned}$$

Dostaneme tedy

$$\begin{aligned} b_1 &= \frac{-162,499}{7,471} = -21,751, \\ b_0 &= 111,111 - (-21,751) \cdot 1,511 = 143,987, \end{aligned}$$

kteréžto odhady jsme už v 2. odstavci použili. Dále můžeme spočítat reziduální součet čtverců. Podle (7) dostaneme

$$S(b_0, b_1) = 9\,595,333 - (-21,751)^2/7,471 = 6\,050,853.$$

Odhad rozptylu (*reziduální rozptyl*) je tedy $S^2 = 6\,050,853/52 = 116,4$. Koeficient determinace vyjde

$$R^2 = 1 - \frac{6050,853}{9585,333} = 36,87 \, \%.$$

Školním prospěchem tedy u dívek vysvětlíme více než třetinu variability hodnot IQ.

Při rozhodování o průkaznosti závislosti IQ na školním prospěchu u dívek použijeme t -statistiku

$$T = \frac{-21,751}{\sqrt{116,4}} \sqrt{7,471} = -5,511. \quad (13)$$

Při rozhodování na běžné 5% hladině porovnáme $|-5,511|$ s kritickou hodnotou $t_{52}(0,05) = 2,007$. Statistika v absolutní hodnotě překračuje zmíněnou kritickou hodnotu, takže testovanou hypotézu $\beta_1 = 0$ musíme zamítnout. Závislost je na 5% hladině statisticky průkazná.

V dnešní době počítačů je k dispozici řada programů, které provedou potřebné výpočty a uvedou výsledky. Například volně šiřitelný program R uvede odhady parametrů β_0, β_1 ve tvaru

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	143.987	6.143	23.439	< 2e-16 ***
zn7	-21.751	3.947	-5.511	1.12e-06 ***

Sloupeček **Std. Error** udává střední chyby obou odhadů. Při prokazování závislosti důležitá t -statistika z (13) je uvedena ve druhém řádku a v předposledním sloupci. Snadno se lze přesvědčit že obě statistiky uvedené ve sloupci **t value** opravdu dostaneme vydělením statistik v prvním a druhém sloupci (viz (11)).

Jak již víme, na 5% hladině jsme závislost IQ na školním prospěchu v sedmé třídě prokázali, protože jsem zamítli hypotézu, podle které by směrnice β_1 popisující skutečnou závislost byla nulová. Interval spolehlivosti pro β_1 spočítaný podle (12) vyjde $(-29,670; -13,832)$. To znamená, že když například porovnáváme střední hodnoty IQ i dívek s průměrnou známkou 1 resp. 2, s 95% pravděpodobností očekáváme rozdíl těchto středních hodnot někde mezi $-29,7$ a $-13,8$.

6. p -hodnota

Poslední sloupeček zmíněného výstupu udává tzv. p -hodnoty. Protože jde o pojem velice často používaný, věnujme se mu trochu podrobněji. Obecně p -hodnota udává pravděpodobnost, že **za platnosti hypotézy** nám mohla vyjít rozhodovací statistika ještě méně příznivá (nebo stejně příznivá) pro testovanou nulovou hypotézu, než jaká nám vyšla v našem testu. Je to nejmenší α z možných hladin, při nichž bychom nulovou hypotézu zamítli.

Konkrétně v našem příkladu při testování hypotézy $H_0 : \beta_1 = 0$ je p -hodnota rovna pravděpodobnosti, že náhodná veličina, která má Studentovo t -rozdělení s 52 stupni volnosti, v absolutní hodnotě dosáhne aspoň hodnoty 5,511. Tato pravděpodobnost je nepatrná, v tabulce je uvedena její hodnota $1,12 \cdot 10^{-6}$.

Použití p -hodnoty při testování hypotéz je velice prosté. Stačí porovnat ji s **předem zvolenou** hladinou α . Hypotézu zamítáme právě tehdy, když je $p \leq \alpha$. Tabulky kritických hodnot pak nepotřebujeme.

7. Odhad v obecnějším modelu

Uvažujme nyní obecnější situaci. Nechť $Y_1, \dots, Y_n$ jsou nezávislé náhodné veličiny takové, že pro $1 \leq i \leq n$ platí

$$EY_i = \sum_{j=1}^k x_{ij}\beta_j,$$

kde x_{ij} jsou známé konstanty, $\beta_1, \dots, \beta_k$ neznámé parametry ($k < n$). Celkem n středních hodnot EY_i tentokrát vyjadřujeme pomocí menšího počtu k neznámých parametrů. Budeme předpokládat, že matice $\mathbf{X} = (x_{ij})$ o n řádcích a k sloupcích má lineárně nezávislé sloupce. Víme tedy, že vektor středních hodnot $E\mathbf{Y}$ je jakousi neznámou lineární kombinací sloupců matice $\mathbf{X}$. Jako speciální případ sem patří nejen regresní přímka zavedená v (1), ale také náhodný výběr z normálního rozdělení, kdy mají všechny náhodné veličiny Y_i stejnou střední hodnotu (jediný parametr, tedy $k = 1$) a konstanty x_i jsou vesměs rovny jedničce.

Zobecněním pojmu rozptyl (variance) náhodné veličiny na náhodný vektor je varianční matice náhodného vektoru. Tato matice má na diagonále rozptyly jednotlivých složek vektoru, mimo diagonálu jsou tzv. kovariance dvojic náhodných veličin. Jsou-li veličiny nezávislé, musí být kovariance nulové, takže varianční matice je pak diagonální. Předpokládejme, že náhodné veličiny $Y_1, \dots, Y_n$ mají všechny stejný rozptyl σ^2 , což je vedle $\beta_1, \dots, \beta_k$ další neznámý parametr. Spolu s dříve uvedeným předpokladem nezávislosti náhodných veličin $Y_1, \dots, Y_k$ to znamená, že varianční matice náhodného vektoru $\mathbf{Y} = (Y_1, \dots, Y_n)^T$ je násobkem jednotkové matice (diagonální matice s jedničkami na diagonále), což zapisujeme ve tvaru $\text{var}\mathbf{Y} = \sigma^2\mathbf{I}$.

Věnujme se nyní úloze odhadu vektoru $\boldsymbol{\beta} = (\beta_1, \dots, \beta_k)^T$. Pro funkci

$$S(\boldsymbol{\beta}) = \|\mathbf{Y} - \mathbf{X}\boldsymbol{\beta}\|^2 = \sum_{i=1}^n \left(Y_i - \sum_{j=1}^k x_{ij}\beta_j \right)^2 \quad (14)$$

platí následující tvrzení:

Věta 12. Funkce $\mathcal{S}(\beta)$ je minimální, právě když je $\beta = \mathbf{b}$, kde je

$$\mathbf{b} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{Y}, \quad (15)$$

a platí

$$\mathcal{S}(\mathbf{b}) = \mathbf{Y}^T \mathbf{Y} - \mathbf{b}^T \mathbf{X}^T \mathbf{Y}. \quad (16)$$

Důkaz tvrzení je poměrně jednoduchý. Funkci $\mathcal{S}(\beta)$ můžeme postupně upravit

$$\begin{aligned} \mathcal{S}(\beta) &= \|\mathbf{Y} - \mathbf{X}\beta\|^2 = \|(\mathbf{Y} - \mathbf{X}\mathbf{b}) + (\mathbf{X}\mathbf{b} - \mathbf{X}\beta)\|^2 \\ &= \|\mathbf{Y} - \mathbf{X}\mathbf{b}\|^2 + \|\mathbf{X}(\mathbf{b} - \beta)\|^2 + 2(\mathbf{b} - \beta)^T \mathbf{X}^T (\mathbf{Y} - \mathbf{X}\mathbf{b}) \end{aligned} \quad (17)$$

$$= \|\mathbf{Y} - \mathbf{X}\mathbf{b}\|^2 + \|\mathbf{X}(\mathbf{b} - \beta)\|^2, \quad (18)$$

když jsme v (17) využili toho, že z výrazu pro $\mathbf{b}$ v (15) plyne $\mathbf{X}^T \mathbf{Y} = \mathbf{X}^T \mathbf{X}\mathbf{b}$ a tudíž třetí člen na pravé straně je nulový. Dále druhý člen v (18) je čtvercem délky vektoru. Je tedy nezáporný, nulový je pouze pro nulový vektor. Vzhledem z předpokládané lineární nezávislosti sloupců matice $\mathbf{X}$ dostaneme nulový vektor pouze v případě, že je $\beta = \mathbf{b}$. Minimální hodnota funkce $\mathcal{S}$ je tedy

$$\begin{aligned} \mathcal{S}(\mathbf{b}) &= \|\mathbf{Y} - \mathbf{X}\mathbf{b}\|^2 = (\mathbf{Y} - \mathbf{X}\mathbf{b})^T (\mathbf{Y} - \mathbf{X}\mathbf{b}) \\ &= \mathbf{Y}^T \mathbf{Y} - 2\mathbf{Y}^T \mathbf{X}\mathbf{b} + \mathbf{b}^T \mathbf{X}^T \mathbf{X}\mathbf{b} \\ &= \mathbf{Y}^T \mathbf{Y} - \mathbf{b}^T \mathbf{X}^T \mathbf{Y}, \end{aligned} \quad (19)$$

když jsme při úpravě (19) vhodně dosadili za $\mathbf{b}$ podle (15).

Víme, že střední hodnota $\mathbf{E}\mathbf{Y}$ náleží k vektorům tvaru $\mathbf{X}\beta$. Tvrzení věty umožňuje určit ten z nich, který je k realizovanému náhodnému vektoru $\mathbf{Y}$ nejbližší.

Minimální hodnota $\mathcal{S}(\mathbf{b})$ funkce $\mathcal{S}$ se nazývá *reziduální součet čtverců*. Nestranným odhadem σ^2 je *reziduální rozptyl*

$$S^2 = \frac{\mathcal{S}(\mathbf{b})}{n - k},$$

kde, jak víme, k je počet sloupců matice $\mathbf{X}$, o nichž stále předpokládáme, že jsou lineárně nezávislé.

Předpokládejme nyní, že náhodné veličiny $Y_1, \dots, Y_k$ mají normální rozdělení. Podobně jako jsme usoudili na normální rozdělení odhadu b_1 , lze dokázat, že odhad $\mathbf{b}$ má normální rozdělení se střední hodnotou $\mathbf{b}$ a varianční maticí $\sigma^2 \mathbf{V}$, kde jsme použili označení $\mathbf{V} = (\mathbf{X}^T \mathbf{X})^{-1}$. Střední chyby odhadů b_j jednotlivých složek vektoru β jsou dány vztahem $\text{S.E.}(b_j) = S\sqrt{v_{jj}}$, kde v_{jj} je j -tý diagonální prvek matice $\mathbf{V}$.

8. Má smysl použít prospěch ze dvou ročníků?

Na našem příkladu si ukažme vyšetřování závislosti na několika regresorech. Kromě známkového průměru z pololetí v sedmé třídě použijeme ještě stejné informace z pololetí osmé třídy. Model rozšíříme na

$$Y_i = \beta_0 + \beta_1 x_i + \beta_2 z_i + E_i,$$

kde z_i jsou zjištěné školní průměry v 8. třídě. Program R dá odhady

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	140.712	6.797	20.703	< 2e-16 ***
zn7	-29.046	7.638	-3.803	0.000384 ***
zn8	9.299	8.343	1.115	0.270256

Koeficient determinace vzrostl z 36,87 % jen na 38,38 %. Uvedené odhady musíme interpretovat velice pečlivě. Odhad $b_1 = -29,046$ lze interpretovat tak, že je to odhad rozdílu IQ pro dívky, které se v 7. třídě lišily o jednotku (například průměr 1 a průměr 2) a v 8. třídě mají stejný průměr. Podobně je $b_2 = 9,299$ odhad rozdílu IQ dívek, které v 7. třídě měly stejný průměr a v 8. třídě se průměrem lišily o jednotku. Chápu, že je tato interpretace obtížná, když nám vychází při horším prospěchu v 8. třídě poněkud větší odhad IQ. Ovšem pozor, na 5% hladině nemůžeme zamítnout hypotézu $\beta_2 = 0$, p -hodnota činí 27 %. V zájmu co nejjednoduššího modelu tedy zůstaneme u hypotézy $\beta_2 = 0$. Není průkazné, že by přidání informace o průměrném prospěchu v 8. třídě skutečně přidalo informaci o hodnotě IQ. Vystačíme tedy s jednodušším modelem, se závislostí na průměru v 7. třídě.

Projevila se také tzv. *multikolinearita*. Obě nezávisle proměnné jsou spolu silně závislé (rok od roku se u jednotlivců průměrný prospěch příliš nemění, v 8. třídě bývá podobný jako v 7. třídě). To mimo jiné způsobilo, že se přesnost odhadů koeficientu β_1 podstatně zhoršila, střední chyba vzrostla téměř na dvojnásobek (z 3,947 u závislosti IQ na prospěchu v 7. třídě na 7,638 u závislosti na prospěchu v 7. i v 8. třídě).

9. Souvisí předpověď IQ podle školního prospěchu také s pohlavím?

Až dosud jsme se zabývali studovanou závislostí pouze u dívek. Porovnání poskytuje obrázek 3, odkud je například zřejmé, že prospěch horší než 2,3 mají pouze hoši. Navíc přímka odhadující závislost u děvčat je poněkud strmější. Proto nepřekvapí, když u chlapců dostaneme poněkud jiné odhady:

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	140.680	4.122	34.125	< 2e-16 ***
zn7	-16.872	1.966	-8.581	9.97e-12 ***

Co může říci o porovnání závislosti u chlapců a dívek statistika?

Začneme nejpodrobnějším modelem (pro každé pohlaví obecně jiná přímka) a pomůžeme si tzv. umělou proměnnou, kterou symbolicky nazveme `hoch`. Pro dívku to bude nula, pro chlapce jednička. Symbolicky pak zapíšeme uvažovanou závislost jako

$$iq = \beta_0 + \beta_1 \cdot zn7 + \beta_2 \cdot hoch + \beta_3 \cdot zn7 \cdot hoch.$$

Dosadíme-li za umělou proměnnou `hoch` hodnotu nula, dostaneme pro dívky

$$iq = \beta_0 + \beta_1 \cdot zn7,$$

kdežto po dosazení jedničky pro chlapce

$$iq = (\beta_0 + \beta_2) + (\beta_1 + \beta_3) \cdot zn7.$$

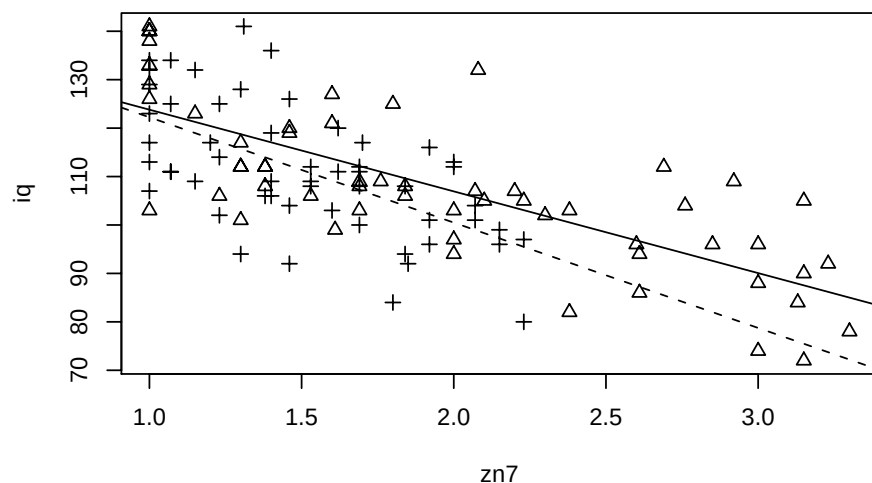
Parametr β_3 tedy opravuje „dívčí“ směrnici přímky na „chlapeckou“, podobně β_2 opravuje „dívčí“ absolutní člen přímky na „chlapecký“. Prostředí R dá odhady

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	143.987	6.125	23.508	< 2e-16 ***
zn7	-21.751	3.935	-5.528	2.31e-07 ***
hoch	-3.308	7.390	-0.448	0.655
zn7:hoch	4.879	4.401	1.108	0.270

Bezprostředním porovnáním snadno ověříme, že u dívek jsme dostali stejné odhady, jako když jsme je vyšetřovali samotné. U chlapců stačí k ověření drobný výpočet; $143,987 + (-3,308) = 140,679$ je až na zaokrouhlení totéž, jako absolutní člen (140,680) v tabulce odhadů pro chlapce, podobně pro směrnici: $-21,751 + 4,879 = -16,872$.

Co když jsou ve skutečnosti přímky pro chlapce a pro dívky rovnoběžné? Takovou hypotézu můžeme v našem podrobném modelu vyjádřit jako požadavek $\beta_3 = 0$ (směrnice obou přímek jsou totožné, „dívčí“ směrnici nemusíme upravovat na „chlapeckou“). Poslední tabulka ukazuje okamžitě i rozhodnutí o hypotéze $\beta_3 = 0$. Protože je příslušná p -hodnota 27 % větší než běžně



Obrázek 3. Souvislost IQ s průměrným prospěchem v 7. třídě u chlapců (soustředěná čára, trojúhelníky) a u děvčat (čárkovaně, křížky).

používaná hladina $\alpha = 5\%$, hypotézu nemůžeme zamítnout a nezbývá, než platnost hypotézy připustit.

Model s rovnoběžnými přímkami dostaneme zjednodušením posledního modelu, vynecháme člen s koeficientem β_3 . Odhady však musíme spočítat znovu. Dostaneme

Coefficients:

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	138.093	3.043	45.376	<2e-16 ***
zn7	-17.852	1.765	-10.116	<2e-16 ***
hoch	4.513	2.198	2.054	0.0424 *

Také tyto odhady je třeba interpretovat velice pečlivě. Kdybychom odhadovali rozdíl mezi IQ u dětí se školním průměrem 1,5 a 2,5, pak můžeme použít odhad $b_1 = -17,852$ jen v případě, že porovnáváme chlapce s chlapcem nebo dívku s dívkou. Jen tehdy nehraje žádnou roli odhad $b_2 = 4,513$. Na druhé straně, když budeme porovnávat předpověď IQ pro chlapce a pro dívku, kteří mají **stejný průměr známek** (aby nehrál roli odhad b_1), pak bude u chlapce odhad větší právě o hodnotu 4,513. Právě o tolik jsou ve svislém směru vůči

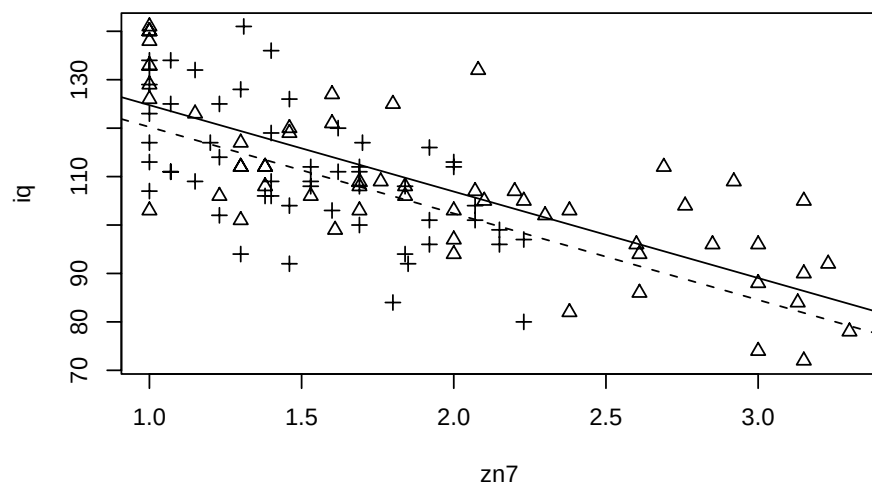
sobě posunuty odhady regresních přímek na obrázku 4. V posledním sloupci uvedená p -hodnota $p = 4,24 \%$ sděluje, že tento rozdíl mezi chlapci a děvčaty je na 5% hladině statisticky průkazný. Znamená to mimo jiné, že už nemůžeme model dále zjednodušit na jedinou přímku pro obě pohlaví.

Spočítejme ještě jeden model. Zapomeňme na školní známky a pokusme se zjistit, zda se co do IQ chlapci a děvčata vůbec liší. K tomu stačí poslední model zjednodušit ještě vyloučením proměnné `zn7`, tedy použít model

$$iq = \beta_0 + \beta_2 \cdot hoch. \quad (20)$$

Vzpomeneme-li si, že modelujeme střední hodnoty a dosadíme-li za proměnnou `hoch` nuly a jedničky, vidíme, že střední hodnota IQ u dívek je β_0 , kdežto u chlapců $\beta_0 + \beta_2$. Parametr β_2 tedy vyjadřuje rozdíl středních hodnot. Říká, o kolik je populační průměr IQ u chlapců větší než u dívek. Výsledné odhady jsou

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	111.111	2.035	54.593	<2e-16 ***
hoch	-3.637	2.840	-1.281	0.203



Obrázek 4. Souvislost IQ s průměrným prospěchem v 7. třídě u chlapců (soustředěná čára, trojúhelníky) a u děvčat (čárkovaně, křížky), když předpokládáme rovnoběžné přímky.

U děvčat je průměrná hodnota IQ rovna $b_0 = 111,111$, kdežto u chlapců je tento průměr roven $b_0 + b_2 = 111,111 + (-3,637) = 107,474$, je tedy právě o 3,637 menší. Není však statisticky průkazné, že by rozdíl mezi středními hodnotami IQ u dívek a chlapců (parametr β_2) byl nenulový. Například interval spolehlivosti pro β_2 vyjde $(-9,266; 1,992)$, takže nula je mezi „přijatelnými“ hodnotami tohoto parametru. Neprokazali jsme (na zvolené 5% hladině významnosti) průkazný rozdíl mezi IQ u chlapců a děvčat.

Zdánlivě je tu neshoda. Při stejném školním průměru očekáváme u chlapců (statisticky prokazatelně) větší hodnotu IQ než u dívek, ale když ke známám nepřihlídneme, situace se změní. Rozdíl už sice není průkazný, ale u děvčat je průměrná hodnota IQ větší. Lze očekávat, že kdybychom použili rozsáhlejší data, rozdíl mezi chlapci a děvčaty (bez přihlížení ke školnímu prospěchu) by byl také průkazný.

Vysvětlení neshody spočívá v tom, že se chlapci od děvčat liší velmi výrazně svým školním prospěchem. Jak jsme již poznamenali, velmi špatné známky mají pouze chlapci. Ty pak průměr IQ u chlapců výrazně snižují. Další pohled na vzájemnou souvislost mezi školními známkami a hodnotami IQ by vneslo vyšetřování závislosti známek na IQ, jak závisí populační průměr známek při dané hodnotě IQ na této hodnotě.

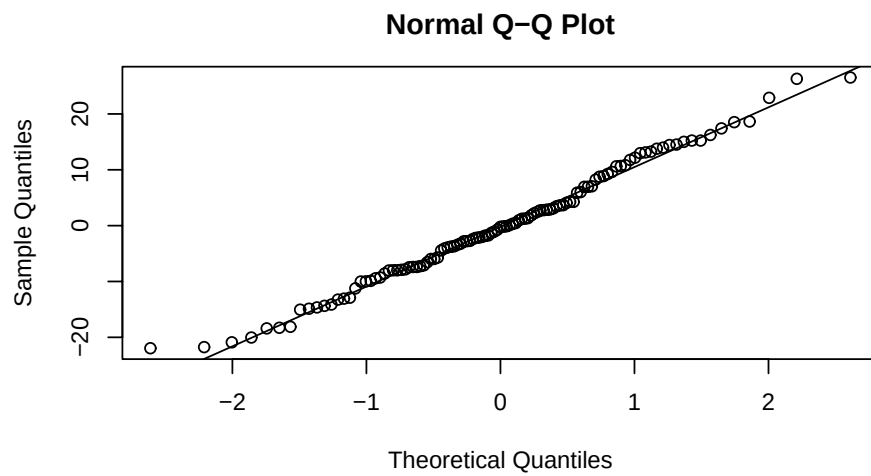
Ještě důležitá poznámka. K rozhodování o shodě středních hodnot u dvou nezávislých výběrů z normálního rozdělení se stejným rozptylem se používá tzv. dvouvýběrový t -test. Model (20) je jen jiným zápisem úlohy, která vede na dvouvýběrový t -test. Hodnotu testové statistiky dvouvýběrového t -testu najdeme v poslední tabulce, je to $T = 1,281$.

10. Závěr

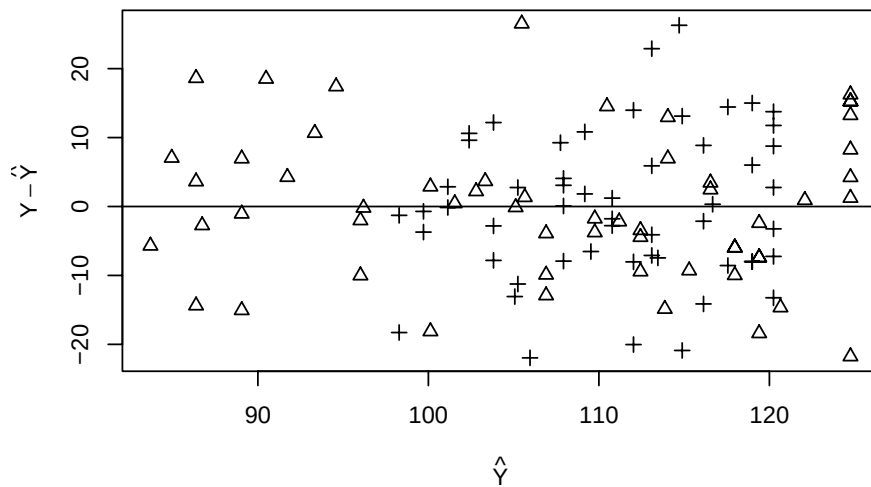
Tato kapitola nese hrdý název regrese, ale jen něco málo naznačuje z toho, co tato snad nejpoužívanější statistická metoda zahrnuje.

Především jsme se věnovali jen modelu, kdy je střední hodnota vysvětlované veličiny *lineární* funkcí neznámých parametrů. Existuje také *regrese nelineární*, kde se toto omezení již neuplatňuje.

Všude v příkladech i u mnoha popisovaných vlastností jsme předpokládali normální rozdělení, což také nemusí být ani přibližně splněno. Někdy si můžeme pomoci vhodnou transformací, kdy například hodnotíme závislost **logaritmu** hmotnosti sušiny kořenové části rostliny v závislosti na procentu cukru v živném roztoku. Když se však pokusíme zjistit, zda množství vykouřených cigaret ovlivňuje výskyt rakoviny plic, bude vysvětlovaná veličina pouze dvouhodnotová. Při vhodně posbíraných datech jde pak o úlohu *logistické regrese*.



Obrázek 5. Regresní diagnostika: ověřování normálního rozdělení.



Obrázek 6. Regresní diagnostika: ověřování tvaru závislosti a konstantního rozptylu (chlapci – trojúhelník, dívky – křížek).

Když už použijeme předpoklad normálního rozdělení, zvolíme model pro střední hodnotu lineární v neznámých parametrech, předpokládáme stejný rozptyl pro všechna pozorování vysvětlované veličiny a předpokládáme nezávislost jednotlivých pozorování této veličiny, měli bychom být schopni aspoň do nějaké míry ověřit, zda všechny tyto předpoklady byly splněny. Tím vším se zabývá *regresní diagnostika*, která s úspěchem využívá různá grafická vyjádření. Ani jedna z ukázek na obrázcích 5 a 6 neukazuje na problém.

První z těchto obrázků naznačuje, že předpoklad normálního rozdělení byl přiměřený datům protože jednotlivé body leží blízko diagonální přímky, není tu patrné nějaké systematické odchylování. Na obrázku 6 jsou body obou pohlaví v podstatě symetricky rozmístěny kolem vodorovné přímky v úrovni nuly, přičemž variabilita kolem této úrovně je přibližně stejná pro malé vyhlazené hodnoty na levé straně a pro velké vyhlazené hodnoty na pravé straně.

Při samotné interpretaci výsledků nesmíme zapomenout na to, jak byla data pořízena, kterou populaci mohou reprezentovat, na jakou populaci můžeme svoje závěry zobecnit.

Literatura

- [1] Likeš J. a Laga J. (1978). *Základní statistické tabulky*, SNTL, Praha.
- [2] Sinclair J. (1826). *An Analysis of the Statistical Account of Scotland*.
- [3] Zvára K. a Štěpán J. (2003). *Pravděpodobnost a matematická statistika*. MATFYZPRESS, Praha.

DEMOGRAFIE A POJISTNÁ MATEMATIKA

Tomáš Cipra

Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín

tomas.cipra@mff.cuni.cz

1. Úmrtnostní tabulky

Jedním z nejdůležitějších výpočetních nástrojů, které se používají v současné demografii jsou *úmrtnostní tabulky* konstruované pomocí demografických metod na základě pozorování rozsáhlých populačních souborů. Úmrtnostní tabulky pro Českou republiku publikuje každoročně na základě nejaktuálnějších úmrtnostních statistik Český statistický úřad.

Jednotlivé sloupce úmrtnostní tabulky obsahují mimo jiné následující hodnoty (viz např. vybrané údaje z úmrtnostní tabulky mužů v České republice za rok 1995 uvedené v tabulce 1):

- *Věk x :*
 - ◇ např. $x = 0, 1, \dots, 103$ (poslední věková skupina znamená „ve věku 103 let a více“);
 - ◇ v pojišťovnách je často $x = 15, 16, \dots$ anebo $x = 18, 19, \dots$
- *Pravděpodobnost úmrtí ve věku x : q_x*
 - ◇ je pravděpodobnost toho, že jedinec, který je naživu ve věku x , zemře před dosažením věku $x + 1$. Např. údaj $q_{50} = 0,008090$ pro muže z tabulky 1 znamená, že z tisíce 50letých mužů jich v průměru osm zemře před dosažením věku 51 let.
- *Pravděpodobnost dožití ve věku x : p_x*
 - ◇ je pravděpodobnost toho, že jedinec, který je naživu ve věku x , se dožije věku $x + 1$. Zřejmě platí $p_x = 1 - q_x$.
- *Počet dožívajících se věku x : l_x*
 - ◇ je počet jedinců z populace l_0 současně narozených jedinců, kteří se dožijí věku x . Např. údaj $l_{50} = 91\,082$ pro muže z tabulky 1 znamená, že z $l_0 = 100\,000$ narozených jedinců mužského pohlaví se jich v průměru 91 082 dožije věku 50 let.

Klíčová slova: Demografie, pojistná matematika.

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839.

- ◇ posloupnost $l_0 \geq l_1 \geq l_2 \geq \dots$ se nazývá dekrementní řád vymírání populace: představuje hypotetický snímek života zvoleného počtu jedinců l_0 , tak jak by vypadal, kdyby se udrželo úmrtnostní chování populace z uvažovaného období (např. z roku 1995 pro tabulku 1);
 - ◇ l_0 se nazývá kořen (radix) úmrtnostní tabulky a obvykle se volí $l_0 = 100\,000$.
 - *Počet zemřelých ve věku x : d_x*
 - ◇ jedná se o počet jedinců z l_0 , kteří zemřou v dokončeném věku x . Např. údaj $d_{50} = 737$ pro muže z tabulky 1 znamená, že z $l_0 = 100\,000$ narozených jedinců mužského pohlaví jich v průměru 737 zemře právě ve věku 50 let.
 - *Střední délka života ve věku x : e_x*
 - ◇ je průměrný počet let, kterých se ještě dožije jedinec ve věku x . Např. údaj $e_{50} = 23,46$ pro muže z tabulky 1 znamená, že 50letý muž může v průměru počítat ještě s 23,46 roky života. Speciálně e_0 je střední délka života (např. podle tabulky 1 byla v roce 1995 v České republice střední délka života mužů 69,96 roku).
- Mezi předchozími veličinami platí řada vztahů, např.

$$d_x = l_x - l_{x+1}, \quad (1)$$

$$q_x = \frac{d_x}{l_x} = \frac{l_x - l_{x+1}}{l_x}, \quad p_x = \frac{l_{x+1}}{l_x}. \quad (2)$$

Úmrtnost žen je prakticky celosvětově a ve všech věkových kategoriích nižší než úmrtnost mužů (a to v některých věkových kategoriích velmi podstatně, např. v České republice v roce 1995 bylo pro muže $q_{50} = 0,008\,090$, zatímco pro ženy jen $q_{50} = 0,003\,459$). Při použití úmrtnostních tabulek v životním pojištění se tato skutečnost zohledňuje tím způsobem, že sazby pojistného pro muže se konstruuji pomocí mužských úmrtnostních tabulek, zatímco sazby pro ženy vytvoří posunutím sazeb pro muže většinou o 5 let (např. žena uzavírající životní pojištění ve věku 40 let tak platí stejné pojistné jako muž uzavírající tentýž typ pojištění ve věku 35 let).

X	q_x	p_x	l_x	d_x	e_x
0	0,007 307	0,992 693	100 000	731	69,96
...	...	...	...	...	...
30	0,001 268	0,998 732	97 280	123	41,52
...	...	...	...	...	...
40	0,002 832	0,997 168	95 526	271	32,18
...	...	...	...	...	...
50	0,008 090	0,991 910	91 082	737	23,46
...	...	...	...	...	...
60	0,019 847	0,980 153	79 982	1 587	15,94
...	...	...	...	...	...

Tabulka 1.

2. Pojištění jako ochrana proti rizikům

Pojišťovnictví jako jedna z klíčových oblastí hospodářství má dvě stránky:

- *etická stránka*: projevuje se v solidaritě ostatních pojištěných s postiženým (tzv. *princip solidarity*);
- *výdělečná stránka*: jedná se o prosperující odvětví pro podnikání (to platí především o životních pojišťovnách, zisky plynoucí z oblasti pojištění majetku začínají v posledních letech celosvětově klesat).

Pojistné riziko je potenciální možnost vzniku *pojistné události*, při níž pojišťovna podle sjednané *pojistné smlouvy* vyplácí *pojistné plnění*. Přitom předmětem pojištění jsou pouze tzv. *čistá rizika* prokazatelně náhodného charakteru (např. doba života, úraz, požár, dopravní havárie apod. na rozdíl od uměle vytvářených *spekulativních rizik*, jako je např. sázková činnost). Nahodilost pojistné události může být *absolutní* (např. požár) nebo *relativní* (např. úmrtí: určité nastane, ale náhodný je jeho okamžik).

Na pojištění lze pohlížet jako na ochranu proti pojistným rizikům: *pojištěný* přenesl svá rizika, jejichž potenciální škodní důsledky jsou z jeho individuálního hlediska neúnosné, na *pojistitele* (*pojišťovnu*), který při dostatečně velkém souboru rizik podobného charakteru (soubor pojistných smluv podobného typu se označuje jako *pojistný kmen*) je schopen celkově převzít rizika s využitím inkasovaného *pojistného* nejen zvládat, ale učinit je předmětem výnosné komerční činnosti. Pojištění tedy v tomto smyslu slouží jako nástroj finanční eliminace negativních důsledků nahodilosti.

Typy pojistných rizik:

- *objektivní riziko*: je dáno objektivními faktory, jako je např. věk, pohlaví, zdravotní stav, profese, charakteristiky pojištěného předmětu a prostředí apod.;
- *subjektivní riziko*: je dáno subjektivními faktory, jako je např. snaha pojištěného zachovat své zdraví a život, vyhnout se střetu se zákonem, zachovat pojištěný předmět ve funkčním stavu apod.;
- *morální riziko*: nastává v situaci, kdy pojištěný nepreferuje jednoznačně zábrannou činnost před vznikem škody (figuruje často v souvislosti s pokusy o pojistné podvody);
- *osobní riziko*: např. riziko předčasné smrti, tělesného poškození, sociální nedostatečnosti při dožití určitého věku apod.;
- *živelní riziko*: riziko přímých škod na majetku v důsledku živelních událostí (např. požáru, povodně apod.);
- *dopravní riziko*: riziko škod vzniklých v souvislosti s dopravním prostředkem (tzv. havarijní či kaskopojištění) nebo s přepravovaným zbožím (tzv. kargopojištění);
- *riziko odcizení a vandalství*: jako podmínka pojistného plnění se často klade překonání předepsaných zabezpečovacích opatření při krádeži nebo zjištění pachatele při vandalství;
- *šomážní riziko*: riziko přerušení provozu v důsledku havárie (např. zkažení zmražených potravin, sankce při nedodržení kontraktu) a riziko ušlého zisku (např. prosperující firma po živelní katastrofě musí vyklidit místo konkurenci);
- *strojní riziko*: riziko havárie či poruchy strojního zařízení v důsledku neodborného zacházení, vady materiálu, chybné technologie apod.;
- *odpovědnostní riziko*: riziko škod způsobených v důsledku jednání pojištěného na zdraví a životě jiné osoby nebo na cizím majetku (např. pojištění odpovědnosti za škody způsobené provozem motorového vozidla, za škody způsobené výkonem povolání daňového poradce, za škody způsobené havárií v domácnosti, za vyrobený léčebný přípravek apod.);
- *sociálně-politické riziko*: zahrnuje válečné operace, etnické konflikty, embarga, stávky apod.;
- *obchodně-finanční riziko*: vyplývá ze změn ekonomických podmínek a dodavatelsko-odběratelských vztahů na domácím a zahraničním trhu, speciálně sem spadá úvěrové riziko spočívající v nebezpečí nesplacení dluhů;
- *moderní rizika*: např. atomové riziko, ekologické riziko, riziko spojené s provozem kosmických těles, riziko AIDS aj.

Z hlediska pojistitele se rizika převzatá v rámci pojistného kmene transformují na tzv. *pojistně-technické riziko* pojistitele, které spočívá v potenciálním nebezpečí, že ve skutečnosti nedojde k vyrovnání mezi přijatým pojistným a vyplaceným pojistným plněním. Toto riziko se měří výší variability mezi očekávaným stavem, z něhož vychází výpočet pojistného, a skutečným stavem, který se odrazí ve vyplaceném pojistném plnění. Podstata pojišťovací činnosti je přitom založena na tom, že s růstem velikosti pojistného kmene (tj. s růstem aktuálního počtu uzavřených pojistných smluv) se pojistně-technické riziko zmenšuje.

Příklad 1. Předchozí tvrzení může být demonstrováno s využitím pravděpodobnostního aparátu např. následujícím způsobem. Nechť v určitém pojištění nastává pojistná událost s pravděpodobností 0,01 (tj. v jednom ze stovky případů), přičemž pojistná událost je spojena se škodou 1 mil. Kč. Z individuálního hlediska se zjevně jedná o značně rizikovou záležitost vyžadující pojistnou ochranu. Proveďte rozbor z hlediska pojistně-technického rizika pojistitele.

Řešení: Z pohledu pojistitele je rozhodující výše škody připadající v průměru na jednu pojistnou smlouvu v jeho pojistném kmene tvořeném n pojistnými smlouvami:

Podle pravděpodobnostního počtu je střední (tj. průměrná) výše škody na jednu pojistnou smlouvu

$$1\,000\,000 \times 0,01n/n = 10\,000 \text{ Kč},$$

takže pojistitel jako cenu za poskytnutí této pojistné ochrany vysadí pojistné právě v této výši (takové pojistné inkasované např. v $n = 100$ pojistných smlouvách, tj. celkem $100 \times 10\,000 = 1\,000\,000$ Kč, zřejmě pokryje právě jednu pojistnou událost, která by v jejich rámci vzhledem k uvedené pravděpodobnosti měla nastat).

Pojistně-technické riziko pojistitele, že pojistné 10 000 Kč v rámci pojistného kmene s n pojistnými smlouvami nebude stačit, je přirozené měřit směrodatnou odchylkou výše škody na jednu pojistnou smlouvu, což je v pravděpodobnostním počtu obvyklá míra pro ocenění chyby při použití průměru (tj. při použití střední hodnoty). Tato směrodatná odchylka je v našem případě (s využitím vzorce pro směrodatnou odchylku binomického rozdělení)

$$1\,000\,000 \times \sqrt{\frac{0,01 \times 0,99}{n}} \approx \frac{99\,500}{\sqrt{n}} \text{ Kč}.$$

Pro $n = 100$ je pojistně-technické riziko ještě poměrně velké (pojistné 10 000 Kč podléhá chybě řádově ve výši $99\,500/\sqrt{100} = 9\,950$ Kč), s rostoucím n

ale klesá, takže např. pro $n = 10\,000$ se již redukuje na přijatelnou úroveň (pojistné 10 000 Kč pak podléhá chybě řádově ve výši $99\,500/\sqrt{10\,000} = 995$ Kč).

Klasifikace pojištění může být provedena zhruba následujícím způsobem:

- *komerční (individuální) pojištění*:
 - ◊ *pojištění osob* (to se někdy dále dělí na kapitálové životní pojištění obsahující spořivou rezervotvornou složku vyplácenou např. při dožití sjednaného věku, rizikové životní pojištění kryjící riziko smrti bez spořivé složky a důchodové pojištění chápáné jako komerční produkt umožňující např. zakoupení doživotní renty);
 - ◊ *úrazové pojištění*;
 - ◊ *pojištění majetku (věcné pojištění)*;
 - ◊ *pojištění odpovědnosti za škody* (někdy se úrazové pojištění, pojištění majetku a pojištění odpovědnosti označují souhrnně jako *neživotní pojištění*, zatímco termín *životní pojištění* je vyhrazen pro pojištění osob);
- *sociální pojištění* zabezpečující úhradu dávek pro případ pracovní neschopnosti, která může být dočasná (pak se jedná o *nemocenské pojištění*) nebo trvalá v důsledku věku či invalidity (pak se jedná o *důchodové* či *penzijní (při)pojištění* garantované státem či penzijními fondy);
- *zdravotní pojištění* garantované státem či v individuální smluvní formě.

Z právního hlediska lze provést následující klasifikaci:

- *dobrovolné pojištění*: sjednává se v závislosti na dobrovolném rozhodnutí klienta (formou pojistné smlouvy);
- *povinné pojištění*:
 - ◊ *povinné smluvní pojištění*: právní předpis určuje povinnost sjednat toto pojištění (formou pojistné smlouvy) jako podmínku určité činnosti (např. pojištění odpovědnosti za škodu způsobenou provozem motorového vozidla označované krátce jako *povinné ručení*, pojištění odpovědnosti provozovatelů civilních letadel, pojištění odpovědnosti za škody z výkonu některých povolání, jako jsou advokáti, notáři, lékaři v nestátních zařízeních, lékárníci, daňoví poradci, auditoři aj.);
 - ◊ *zákonné pojištění*: jeho povinnost ukládá zákon, přičemž se nesjednává pojistná smlouva (u nás už jediné zákonné pojištění odpovědnosti organizace za škodu při pracovním úrazu a nemoci z povolání).

3. Pojištění osob

V rámci pojištění osob lze sjednat:

- *pojištění pro případ smrti*: pojistnou událostí je smrt pojištěného;
- *pojištění pro případ dožití*: pojistnou událostí je dožití sjednaného věku pojištěným;
- *smíšené pojištění*: pojistnou událostí je smrt *nebo* dožití sjednaného věku pojištěným;
- *důchodové pojištění*: je v podstatě speciálním případem pojištění pro případ dožití s pravidelně se opakujícím pojistným plněním ve formě výplaty důchodu.

Účastníci pojištění jsou:

- *pojistitel* je provozovatel pojištění (většinou pojišťovna);
- *pojistník* je fyzická nebo právnická osoba, která s pojistitelem uzavřela pojistnou smlouvu a má povinnost platit pojistné;
- *pojištěný (pojištěnec, účastník)* je fyzická osoba, na jejíž život a zdraví je pojištění sjednáno (pojistník a pojištěný mohou být případně tatáž osoba);
- *oprávněná osoba (obmyšlený)* je fyzická nebo právnická osoba, která má právo na výplatu pojistného plnění v případě smrti pojištěného (oprávněnou osobou ale také může být např. banka v pojištění kryjícím riziko smrti dlužníka u hypoték).

Příklad 2. Uvažujte situaci, kdy 40letý muž uzavře na dobu 20 let pojištění na pojistnou částku 500 000 Kč, a to (1) pro případ smrti; (2) pro případ dožití; (3) smíšené pojištění. Co se bude dít při těchto jednotlivých variantách, jestliže

- pojištěný zemře před dožitím věku 60 let?
- pojištěný se dožije věku 60 let?

Řešení:

(1) Pojištění pro případ smrti:

- oprávněným osobám (většinou pozůstalým) vyplatí pojišťovna částku 500 000 Kč;
- pojištění zanikne bez náhrady.

(2) Pojištění pro případ dožití:

- pojištění zanikne bez náhrady;
- pojištěnému vyplatí pojišťovna částku 500 000 Kč.

(3) Smíšené pojištění:

- oprávněným osobám (většinou pozůstalým) vyplatí pojišťovna částku 500 000 Kč;
- pojištěnému vyplatí pojišťovna částku 500 000 Kč.

(2') Jako příklad důchodového pojištění lze uvažovat případ, kdy 40letý muž uzavře pojištění na měsíční důchod 5 000 Kč odložený k věku 60 let. Pak na začátku každého měsíce, kterého se dožije po dosažení věku 60 let, inkasuje tento pojištěný částku 5 000 Kč.

Pojistné může být:

- *jednorázové*: zaplatí se najednou při uzavření smlouvy;
- *běžné*: platí se opakovaně, a to obvykle v pravidelných pojistných obdobích splátkami stejné výše (navíc pojišťovna většinou zohledňuje určitým zvýhodněním menší frekvence placení pojistného, např. 5procentní slevou roční pojistné zaplacené na začátku každého pojistného roku vůči měsíčnímu pojistnému placenému vždy na začátku každého pojistného měsíce);
- *nettopojistné*: je vypočteno tak, aby v průměru pojišťovně krylo pojistná plnění;
- *bruttopojistné*: je nettopojistné rozšířené o složky na pokrytí správních nákladů pojišťovny a případných nepříznivých škodních výchylek formou bezpečnostní přírážky;
- *valorizované*: pojistné se zvyšuje nejčastěji podle momentálního vývoje inflace.

Sazebník pojišťovny uvádí pro jednotlivé pojistné produkty výši pojistného, přičemž v úvahu bere:

- *pohlaví pojištěného*: vzhledem k tomu, že úmrtnost žen je prakticky ve všech věkových kategoriích nižší než úmrtnost mužů (a to pro některé věkové kategorie velmi podstatně), platí žena v pojištění pro případ smrti (resp. v pojištění pro případ dožití) menší pojistné (resp. větší pojistné) než muž stejného věku;
- *vstupní věk pojištěného*: v České republice se stanovuje jednotně jako rozdíl kalendářního roku uzavření pojištění a roku narození pojištěného (např. pojištěný narozený 17. prosince 1960, který uzavřel pojištění dne 17. ledna 2003, má vstupní věk $2003 - 1960 = 43$, i když z čistě matematického hlediska je v okamžiku uzavření pojištění jeho věk jen $42\frac{1}{12}$ a např. v německém pojistném systému by byl jeho vstupní věk vyhodnocen opravdu jako 42);

- *pojistná doba:*

- ◊ *dočasné pojištění:* jeho pojistná doba je předem smluvně omezena (např. pojištění pro případ smrti sjednané ve věku 40 let na dobu 10 let pokrývá riziko smrti jen po tuto smluvní dobu, takže např. při úmrtí pojištěného ve věku 51 let již pojišťovna pozůstalým klienta nic nevyplatí);
- ◊ *trvalé pojištění:* jeho pojistná doba není předem smluvně omezena (např. v pojištění doživotního důchodu probíhá výplata až do smrti klienta bez jakéhokoli deterministického smluvního omezení);
- ◊ *pojištění s odkladem:* povinnost pojistného plnění pojišťovnou je odložena o sjednanou dobu (např. důchod odložený k věku 60 let může být sjednán za jednorázové pojistné zaplacené již ve věku 40 let, i když vlastní výplaty důchodu začnou až od věku 60 let).

Pojištění více životů je pojištění, u něhož pojistné plnění závisí na životě či smrti dvou nebo více osob, např. smíšené pojištění dvojic, vdovský a sirotčí důchod aj. (např. vdovský důchod se vyplácí jen v situaci, kdy žena zůstává naživu po smrti muže).

Skupinové pojištění je hromadně sjednávané pojištění pro skupinu osob (např. zaměstnavatel hromadně pojistí své zaměstnance). Vzhledem k administrativnímu zjednodušení (např. hromadné placení pojistného přes jednu mzdovou účtárnu) lze obvykle stanovit nižší pojistné a případně ignorovat rozdílné vstupní věky pojištěných.

Sdružené pojištění je pojištění více rizik v rámci jedné pojistné smlouvy (např. úrazové připojištění nebo tzv. rodinná pojištění sloužící ke komplexní pojistné ochraně rodiny, jako je např. sdružené pojištění mládeže pokrývající rizika smrti a úrazu jak u rodiče (rodičů), tak u dítěte (dětí).

Zajištění je pojištění pojišťovny formou přenesení (*cese*) části rizik pojišťovnou (*prvopojistitelem*, *cedentem*) na jiného pojistitele (*zajistitele*, *cessionáře*) za cenu odstoupení části inkasovaného pojistného (tzv. *zajistné*) prvopojistitelem zajistiteli. Zajišťovny (tj. pojišťovny specializující se na zajištění) působí obvykle v mezinárodním měřítku a mohou se opřít o obrovské pojistné kmeny.

Základní nástroje výpočetních postupů v rámci pojištění osob jsou:

- *dekrementní nástroje:* především úmrtnostní tabulky (viz odstavec 1);
- *finanční nástroje:* jedná se především o tzv. pojistně-technickou úrokovou míru, kterou životní pojišťovna používá pro diskontování např. při kalkulaci pojistného.

Ve všech příkladech uváděných v tomto textu se pracuje s mužskými úmrtnostními tabulkami v České republice pro rok 1995 a s pojistně-technickou úrokovou mírou 4%.

Příklad 3. Jaké pojistné by měla jednorázově požadovat životní pojišťovna od 30letého muže, který s ní uzavírá *pojištění pro případ dožití* věku 50 let s pojistnou částkou 100 000 Kč? (Pojišťovna tedy vyplatí pojištěnému tuto částku, jakmile se pojištěný dožije věku 50 let; při jeho úmrtí před dožitím tohoto věku pojištění zanikne bez náhrady. V praxi se ovšem tento typ pojištění, aby byl pro klienty přijatelnější, obvykle nabízí s tzv. výhradou vrácení dosud zaplaceného pojistného oprávněné osobě v případě smrti pojištěného.)

Řešení:

Podle principu fiktivního souboru pojišťovna kalkuluje s tím, že odpovídající počet 30letých pojištěných bude $l_{30} = 97\,280$, z nichž se $l_{50} = 91\,082$ dožije věku 50 let (viz tabulka 1). Pojišťovna tedy po uplynutí 20 let bude vyplácet částku

$$100\,000\,l_{50} = 100\,000 \times 91\,082 = 9\,108\,200\,000 \text{ Kč.}$$

se současnou hodnotou

$$9\,108\,200\,000 \times (1/1,04)^{20} = 4\,156\,863\,583 \text{ Kč.}$$

Podle principu ekvivalence však tuto částku musí zaplatit ve formě pojistného 30letí pojištění, takže na každého z nich vyjde

$$4\,156\,863\,583/l_{30} = 4\,156\,863\,583/97\,280 = 42\,731 \text{ Kč.}$$

Jestliže pomineme provozní náklady pojišťovny (tj. omezíme se na nettopojistné místo skutečně požadovaného bruttopojistného) a pomineme případný zisk z vyššího úročení kapitálu než pouhými 4%, pak by příslušné pojistné mělo být 42 731 Kč (stejně pojistné by zřejmě vzhledem k 5letému posunu zmíněnému v odstavci 1 zaplatila žena uzavírající ve věku 25 let pojištění pro případ dožití věku 45 let). Celkový výpočet lze zřejmě shrnout do tvaru

$$\frac{100\,000\,l_{50}\,v^{20}}{l_{30}} = 100\,000\frac{l_{50}\,v^{50}}{l_{30}\,v^{30}} = 100\,000\frac{D_{50}}{D_{30}},$$

kde $v = 1/(1+i)$ je diskontní faktor odpovídající pojistně-technické úrokové míře i a $D_x = l_x v^x$ je příkladem tzv. komutačních čísel (viz dál).

Komutační čísla jsou pomocné hodnoty, které vznikají finančním diskontováním hodnot z úmrtnostních tabulek. Pojišťovny je používají obvykle v tabelované formě pro zjednodušení a zpřehlednění pojistně-matematických výpočtů. Nejdůležitější komutační čísla jsou

- *diskontovaný počet dožívajících se věku x :*

$$D_x = l_x v^x, \quad (3)$$

kde $v = 1/(1+i)$ je diskontní faktor odpovídající pojistně-technické úrokové míře i ;

- *diskontovaný počet zemřelých ve věku x :*

$$C_x = d_x v^{x+1}; \quad (4)$$

- *komutační čísla vyšších řádů:*

$$N_x = D_x + D_{x+1} + \dots, \quad M_x = C_x + C_{x+1} + \dots \quad (5)$$

Ve vzorci (5) se na rozdíl od (4) diskontuje přes $x+1$ období, neboť počet zemřelých d_x odpovídá stavu až na konci daného roku. Přitom např. platí

$$C_x = d_x v^{x+1} = (l_x - l_{x+1})v^{x+1} = D_x v - D_{x+1}.$$

Velice výhodné pro praktický výpočet komutačních čísel jsou tabulkové procesory typu Excel aj. V tabulce 2 jsou uvedena vybraná komutační čísla odpovídající údajům z úmrtnostní tabulky mužů v České republice za rok 1995 (viz tabulka 1).

x	q_x	D_x	C_x	N_x	M_x
0	0,007 307	100 000,0	702,9	2 384 830,5	8 275,7
...	...	...	...	...	...
30	0,001 268	29 993,2	36,5	608 124,5	6 603,8
...	...	...	...	...	...
40	0,002 832	19 897,0	54,1	356 748,1	6 175,9
...	...	...	...	...	...
50	0,008 090	12 816,4	99,7	191 545,9	5 449,2
...	...	...	...	...	...
60	0,019 847	7 603,1	145,1	88 148,1	4 212,8
...	...	...	...	...	...

Tabulka 2.

Např. je

$$D_{30} = l_{30} v^{30} = 97\,280 \times (1/1,04)^{30} = 29\,993,2;$$

$$C_{30} = d_{30} v^{31} = 123 \times (1/1,04)^{31} = 36,5;$$

$$N_{30} = D_{30} + D_{31} + \dots = 29\,993,2 + 28\,803,2 + \dots = 608\,124,5$$

a v příkladě 3 je opravdu

$$100\,000 \frac{D_{50}}{D_{30}} = 100\,000 \frac{12\,816,4}{29\,993,2} = 42\,731 \text{ Kč.}$$

Příklad 4. Jaké je jednorázové a běžné roční pojistné pro 40letého muže v *dočasném pojištění pro případ smrti* na dobu 20 let s pojistnou částkou 100 000 Kč? (Pojišťovna tedy vyplatí tuto částku oprávněné osobě v případě smrti pojištěného; při jeho dožití věku 60 let však pojištění zanikne bez náhrady. Tento typ pojištění se na pojistném trhu často nabízí pod názvem *životní úvěrové pojištění* a jeho sjednání si banky často kladou jako podmínku poskytnutí úvěru: v případě smrti dlužníka, který si musel sjednat toto pojištění s pojistnou částkou ve výši dlužné částky jako podmínku získání úvěru, je příslušná pojistná částka v první řadě použita k umoření dluhu. Někdy se toto pojištění uzavírá s klesající pojistnou částkou za předpokladu, že dlužná částka bude dlužníkem postupně umořována a v případě jeho smrti tedy k umoření bude zbývat podstatně méně, než je původní výše pojistné částky. V každém případě však nutnost sjednání takového pojištění nezanedbatelně prodrazí např. hypoteční úvěr, přičemž při hypotečním úvěru musí být navíc ještě pojištěna příslušným majetkovým pojištěním odpovídající nemovitost.)

Řešení:

Jestliže symbolem P označíme běžné roční pojistné, pak princip ekvivalence

$$\begin{aligned} & \text{současná hodnota očekávaného pojistného} \\ &= \text{současná hodnota očekávaných pojistných nároků} \end{aligned}$$

dává pro dočasné pojištění pro případ smrti se vstupním věkem x , pojistnou dobou n a pojistnou částkou S rovnici tvaru

$$P l_x + P l_{x+1} v + \dots + P l_{x+n-1} v^{n-1} = S d_x v + S d_{x+1} v^2 + \dots + S d_{x+n-1} v^n. \quad (6)$$

Jestliže obě strany rovnice (6) vydělíme hodnotou l_x a vzniklé zlomky rozšíříme hodnotou v^x , pak dostaneme

$$\begin{aligned} & P \frac{l_x v^x + l_{x+1} v^{x+1} + \dots + l_{x+n-1} v^{x+n-1}}{l_x v^x}, \\ & S \frac{d_x v^{x+1} + d_{x+1} v^{x+2} + \dots + d_{x+n-1} v^{x+n}}{l_x v^x}, \end{aligned}$$

neboli po dosažení komutačních čísel typu (3) a (4)

$$P \frac{D_x + D_{x+1} + \dots + D_{x+n-1}}{D_x} = S \frac{C_x + C_{x+1} + \dots + C_{x+n-1}}{D_x}. \quad (7)$$

Tuto rovnici však lze ještě upravit dále s použitím komutačních čísel typu (5) do tvaru

$$P \frac{N_x - N_{x+n}}{D_x} = S \frac{M_x - M_{x+n}}{D_x}. \quad (8)$$

Pravá strana rovnice (8)

$$S \frac{M_x - M_{x+n}}{D_x} \quad (9)$$

pak zřejmě představuje příslušné jednorázové pojistné, zatímco odpovídající běžné roční pojistné P získáme řešením rovnice (8) jako

$$P = S \frac{M_x - M_{x+n}}{N_x - N_{x+n}}. \quad (10)$$

V našem příkladě je konkrétně $x = 40$, $n = 20$, $S = 100\,000$, takže požadované jednorázové pojistné dostaneme dosažením komutačních čísel z tabulky 2 do (9) jako

$$S \frac{M_{40} - M_{60}}{D_{40}} = 100\,000 \frac{6\,175,9 - 4\,212,8}{19\,897,0} = 9\,866 \text{ Kč}$$

a běžné roční pojistné dosažením hodnot z tabulky 2 do (10) jako

$$P = S \frac{M_{40} - M_{60}}{N_{40} - N_{60}} = 100\,000 \frac{6\,175,9 - 4\,212,8}{356\,748,1 - 88\,148,1} = 731 \text{ Kč}.$$

Rovněž 35letá žena by tedy zaplatila za tuto 20letou pojistnou ochranu jednorázově 9 866 Kč nebo ročně (do případného úmrtí) 731 Kč.

4. Úrazové pojištění

Pojistné plnění v rámci úrazového pojištění pojišťovna vyplácí, jestliže došlo k tělesnému poškození (popř. smrti) pojištěného v důsledku úrazu (úrazové pojištění se velmi často uzavírá jako připojištění k jiným pojistným produktům). Konkrétně se přitom jedná o následující možnosti pojistného plnění:

- *pojistné plnění za dobu nezbytného léčení* (DNL) tělesného poškození způsobeného úrazem: pojišťovna vyplatí tolik procent ze sjednané pojistné

částky, kolik odpovídá podle oceňovacích tabulek průměrné době nezbytného léčení příslušného tělesného poškození;

- *pojistné plnění za trvalé následky úrazu (TNÚ)*: pojišťovna vyplatí tolik procent ze sjednané pojistné částky, kolik odpovídá podle oceňovacích tabulek rozsahu trvalých následků úrazu po jejich ustálení;
- *pojistné plnění za smrt úrazem (SÚ)*.

V rámci úrazového pojištění pojišťovny obvykle provádějí důslednou klasifikaci rizika, např. rozlišují činnost pracovní a mimopracovní a různé rizikové skupiny podle povolání, registrovaných sportů apod., např.

- riziková skupina 1: všechny profese mimo 2, 3, 4; pěší turistika, ...;
- riziková skupina 2: manuální dělnické profese mimo 3; sporty mimo 1, 3, 4;
- riziková skupina 3: pokrývač, dřevorubec, policista, ...; cyklistika, vzpírání, ...;
- riziková skupina 4: kaskadér, rizikový zpravodaj, ...; box, horolezectví, parašutismus, ...

Příklad 5. Měsíční pojistné pro základní pojistnou částku 1 000 Kč v rizikové skupině 2 předepisuje pojišťovna pro DNL ve výši 2,00 Kč, pro TNÚ ve výši 0,40 Kč a pro SÚ ve výši 0,20 Kč. Určete roční pojistné, jestliže ve skutečnosti byly sjednány pojistné částky 70 000 Kč pro DNL, 300 000 Kč pro TNÚ a 300 000 Kč pro SÚ.

Řešení:

Měsíční pojistné je zřejmé:

$$70 \times 2,00 + 300 \times 0,40 + 300 \times 0,20 = 320 \text{ Kč.}$$

Jestliže je pojištěný ochoten zaplatit za celý rok dopředu formou ročního pojistného, poskytuje pojišťovna obvykle slevu, takže při 5procentní slevě je výsledné roční pojistné:

$$12 \times 320 \times 0,95 = 3\,648 \text{ Kč.}$$

5. Pojištění majetku

Pojištění majetku umožňuje sjednat následující pojistnou ochranu:

- *Pojištění majetku obyvatelstva*, např.:
 - ◊ *pojištění domácnosti*: předmětem pojištění je soubor zařízení domácnosti, přičemž se obvykle jedná o sdružené krytí různých rizik (většina živelních rizik, vodovodní riziko, riziko odcizení při překonání zabezpečujících překážek aj.);

- ◇ *pojištění budov (staveb)*: předmětem pojištění je budova (i ve výstavbě), např. rodinný dům, nájemní obytný dům, rekreační stavba, hospodářská budova, drobná stavba (garáž) apod., přičemž se obvykle jedná o sdružené krytí různých rizik (živelní rizika, vodovodní riziko, náraz dopravních prostředků, riziko odcizení (stavebních součástí), odpovědnost vyplývající z vlastnictví budovy aj.);
- ◇ *havarijní pojištění (kaskopojištění)*: předmětem pojištění jsou škody na motorových vozidlech, přičemž se obvykle jedná o sdružené krytí různých rizik (živelní rizika, riziko havárie formou střetu, nárazu apod., riziko odcizení často podmíněné instalací zabezpečujících zařízení na vozidle apod.). Pojistné plnění se většinou provádí do výše časové (tj. amortizované) ceny vozidla snížené o cenu upotřebitelných zbytků, a to navíc maximálně do výše sjednané pojistné částky. Pro havarijní pojištění je obvykle typická *spoluúčast* v předem sjednané výši. Pojistné závisí často na typu a stáří vozidla, obsahu motoru, výši zvolené spoluúčasti, regionu (velkoměsto - město - venkov), účelu využívání (služební či rekreační) apod., přičemž se snižuje o tzv. *bonusy* (smluvně zaručené slevy pojistného podle počtu minulých bezškodních let) a případně zvyšuje o tzv. *malusy* (přírážky k pojistnému podle počtu a výše uplatněných pojistných nároků v minulosti).
- *Pojištění průmyslových a podnikatelských rizik*, např.
 - ◇ *živelní pojištění*: je základem pojistné ochrany každého podnikatelského subjektu. Kryje většinu živelních rizik (požár, výbuch, blesk, vichřice, krupobití, povodeň, sesuv lavin, zřícení skal, zemětřesení aj.);
 - ◇ *strojní pojištění*: kryje rizika havárie strojů;
 - ◇ *pojištění pro případ přerušení provozu a ušlého zisku (šomážní pojištění)*: navazuje na živelní a strojní pojištění v tom smyslu, že kryje následné škody. Tyto následné škody často výrazně převyšují přímé škody na majetku. Pojistné plnění se obvykle vyplácí po určitou sjednanou dobu (např. ve výši sumy fixních nákladů a očekávaného zisku v tomto období);
 - ◇ *pojištění proti odcizení*: kryje majetek podnikatelského subjektu pro případ odcizení, poškození nebo zničení majetku jednáním pachatele, které směřovalo ke krádeži vloupáním nebo loupežným přepadením. Pojistné je obvykle diferencováno podle úrovně zabezpečujících opatření;
 - ◇ *dopravní pojištění (pojištění přepravy, kargopojištění)*: kryje riziko zničení, odcizení nebo ztráty věci ve vnitrostátní přepravě nebo v zahraničním obchodě (vzhledem k dlouhým dopravním trasám a rizikovosti při přepravě může cena dopravy převyšovat cenu zboží);

- ◇ *kaskopojištění podnikatelských subjektů*: sem spadá především havarijní pojištění motorových vozidel (např. kamionů), ale také pojištění letectvého či námořního kaska apod.;
- ◇ *pojištění úvěru*: kryje riziko ztrát v případě nesplacení poskytnutého úvěru;
- *Pojištění zemědělských rizik*, kam vedle druhů pojištění uplatňovaných v podnikatelské sféře (viz výše) spadá především:
 - ◇ *pojištění plodin*: kryje majetkové škody na rostlinné produkci, např. krupobitní pojištění, pojištění úrody plodin, pojištění proti vybraným rizikům (např. jarní mráz nebo pojištění vinné révy či chmele);
 - ◇ *pojištění hospodářských zvířat*: vztahuje se na soubory hospodářských zvířat, ale také případně na jednotlivá zvířata (např. závodní koně apod.);
 - ◇ *pojištění lesů*.

Jeden z hlavních rozdílů při výpočtu pojistného v rámci pojištění majetku ve srovnání s pojištěním osob spočívá v tom, že v pojištění majetku se většinou používá tzv. *přirozené pojistné*. Přirozené pojistné je vykalkulované tak, že pokryje pojištěné riziko na jeden rok (či obecněji na jedno pojistné období) dopředu; na konci tohoto roku je pojistné inkasované od kmene pojištěných beze zbytku spotřebováno, neboť odpovídá pravděpodobnosti, že v daném roce nastane pojistná událost (např. havárie či krádež motorového vozidla v havarijním pojištění). Naproti tomu pojistné v pojištění osob se kalkuluje na celou sjednanou pojistnou dobu, která je většinou víceletá (viz odstavec 3).

Základní vzorec pro výpočet pojistného v pojištění majetku je možné zformulovat následujícím způsobem:

$$P = q_1 q_2 S \quad (11)$$

S ... sjednaná pojistná částka

q_1 ... škodní frekvence popisující pravděpodobnost výskytu příslušné pojistné události v pojistném kmeni:

$$q_1 = \frac{\text{odhadnutý počet pojistných událostí v daném roce}}{\text{odhadnutý počet pojistných smluv v daném roce}} \quad (12)$$

q_2 ... škodní rozsah popisující průměrnou výši škody vzhledem k průměrné výši pojistné částky (na rozdíl od pojištění osob totiž v pojištění majetku bývá výše škody a odpovídajícího pojistného plnění obvykle menší než sjednaná pojistná částka):

$$q_2 = \frac{\text{průměrná škoda v jedné pojistné události}}{\text{průměrná pojistná částka v jedné pojistné smlouvě}} \quad (13)$$

Příklad 6. Ze statistik vedených pojišťovnou v rámci pojištění domácnosti byla odhadnuta škodní frekvence ve výši 3,11% (tj. pravděpodobnost škodní události v tomto typu majetkového pojištění je 0,031 1) a škodní rozsah ve výši 12,7% (tj. průměrná výše škody představuje v průměru 12,7% sjednané pojistné částky). Jaké pojistné by mělo být předepsáno pro pojistnou částku 300 000 Kč?

Řešení:

Podle (11) pro $q_1 = 0,031\ 1$, $q_2 = 0,127$ a $S = 300\ 000$ je

$$P = q_1 q_2 S = 0,031\ 1 \times 0,127 \times 300\ 000 = 1\ 185\ \text{Kč}.$$

Príslušné pojistné by tedy mělo být 1 185 Kč. Opět se ovšem jedná pouze o nettopojistné bez zakalkulovaných správních nákladů pojišťovny. Konečné bruttopojistné může být podstatně vyšší.

6. Pojištění odpovědnosti

Pojistnou událostí je zde vznik povinnosti pojištěného nahradit škodu (na životě, zdraví, majetku), pokud tato škoda vznikla v souvislosti s činností nebo vztahem pojištěného blíže specifikovanými v pojistné smlouvě. Pojištění se obvykle nevztahuje na škodu způsobenou úmyslně, na škodu nad rámec stanovený právními předpisy (např. vyplývající z trestné činnosti pojištěného), na škodu, za kterou pojištěný odpovídá přímým příbuzným nebo osobám žijícím s ním ve společné domácnosti, při nesplnění povinnosti k odvrácení škody apod. U nás stále existují tři formy odpovědnostního pojištění:

- *Smluvní pojištění odpovědnosti*, kdy ekonomický subjekt (občan, podnikatelský subjekt) pojišťuje svou odpovědnost za škodu na základě vlastního uvážení, např.:
 - ◊ *pojištění odpovědnosti za škody občana v běžném občanském životě*: toto pojištění zahrnuje např. škody způsobené činností občana v běžném životě (s vyloučením činnosti v rámci pracovně-právních vztahů), provozem domácnosti, jednáním nezletilých dětí, nezávodním provozováním sportů (včetně používání jízdních kol) apod.;
 - ◊ *pojištění odpovědnosti za škody občana - vlastníka, držitele, nájemce nebo správce nemovitostí či vlastníka budovy ve stavbě nebo v demolici*;
 - ◊ *pojištění odpovědnosti za výrobek*.
- *Povinné smluvní pojištění odpovědnosti*, kdy určité ekonomické subjekty jsou povinny na základě právních předpisů sjednat pojistnou smlouvu jako podmínku provozování určité činnosti, např.:

- ◇ *pojištění odpovědnosti za škodu způsobenou provozem motorového vozidla* (tzv. *povinné ručení*);
- ◇ *pojištění odpovědnosti za škody z výkonu povolání* (např. advokáti, notáři, stomatologové, veterinární lékaři, lékaři v nestátních zdravotnických zařízeních, lékárníci, architekti, daňoví poradci, auditoři).
- *Zákonné pojištění odpovědnosti*, kdy vymezeným subjektům je ze zákona stanovena povinnost platit pojistné, aniž by byla sjednána pojistná smlouva. U nás existuje už jen jeden typ zákonného pojištění odpovědnosti, a to:
 - ◇ *zákonné pojištění odpovědnosti organizace za škodu při pracovním úrazu a nemoci z povolání*.

7. Penzijní pojištění

Penzijní pojištění je důchodové pojištění větších populačních skupin, které jsou často určitým způsobem vymezeny: např. obyvatelé celého regionu či státu (*sociální důchodové zabezpečení*), zaměstnanci koncernu či profesního sdružení (penzijní pokladny), účastníci penzijního připojištění se státním příspěvkem (týká se *penzijních fondů* v České republice) apod.

Pro penzijní pojištění různých typů je společné, že za příslušné pojistné (*příspěvky*) se pojišťuje riziko sociální nedostatečnosti a invalidity s pojistným plněním většinou ve formě různých *penzí* (*dávek*):

- *starobní penze* po dosažení stanoveného věku;
- *výsluhová penze* po dosažení stanoveného počtu let účasti v pojištění;
- *invalidní penze* po nastoupení plné (nebo částečné) invalidity;
- *pozůstalostní penze* (především vdovská a sirotčí penze);
- *jednorázové vyrovnání* místo některých z předchozích penzí.

V souvislosti s penzijním pojištěním se mluví o tzv. *penzijních plánech*, které mohou být klasifikovány

- *podle způsobu výpočtu příspěvků a dávek*:
 - ◇ *příspěvkově definovaný*: v takovém penzijním plánu se podle výše zaplacených příspěvků určuje výše dávek;
 - ◇ *dávkově definovaný*: v takovém penzijním plánu se podle výše požadovaných dávek určuje výše příspěvků.
- *podle způsobu financování*:
 - ◇ *fondový (penzijní fond)*: vytváří fondy (rezervy) značných objemů ke krytí svých pozdějších závazků (kdyby došlo k předčasnému rozpuštění takového penzijního fondu, pak díky rezervám by se fond mohl vyrovnat v plné výši se všemi svými závazky, např. by mohl nechat doběhnout všechny priznané starobní penze);

◇ *průběžně financovaný (nefondový, pay-as-you-go)*: příspěvky daného roku pokryjí právě objem dávek v tomto roce. Většina penzijních plánů organizovaných státem je průběžného typu (včetně našeho sociálního důchodového zabezpečení, i když se zde stále důrazněji mluví o nutnosti důchodové reformy a vůbec reformy celého systému veřejných financí).

Penzijní připojištění pomáhá státu čelit rostoucímu *důchodovému břemenu* (míře závislosti populace v poproduktivním věku na populaci v produktivním věku). Stejně jako v celém vyspělém světě i u nás populace postupně stárne v tom smyslu, že ve věkovém složení populace se začínají prosazovat starší ročníky. Penzijní připojištění bývá většinou fondového typu organizované různými penzijními fondy.

Literatura

- [1] Bowers N.L. (1986). *Actuarial Mathematics*. Society of Actuaries, Itasca, první vydání.
- [2] Cipra T. (1999). *Pojistná matematika: Teorie a praxe*. Ekopress, Praha.
- [3] Cipra T. (2004). *Zajištění v pojišťovnictví*. Grada, Praha.

POUŽITÍ PROGRAMU MUPAD VE STŘEDOŠKOLSKÉ VÝUCE

Jaromír Antoch¹, Michal Čihák², Jan Prachař²

¹Univerzita Karlova v Praze, Matematicko-fyzikální fakulta, Katedra pravděpodobnosti a matematické statistiky, Sokolovská 83, 186 75 Praha 8 – Karlín;

²Gymnázium F. M. Pelcla v Rychnově nad Kněžnou
jaromir.antoch@mff.cuni.cz, michal.cihak@mff.cuni.cz

1. Úvod

Nedávno vyšel v časopise MFI velmi zajímavý článek [6], zabývající se využitím moderního CAS systému *MuPAD* ve výuce matematiky na střední škole. Připomeňme, že CAS je zkratka anglického *Computer Algebra Systems*, což v překladu znamená *systémy pro počítačovou algebru*. Již z názvu je patrné, že se jedná o programy umožňující provádět symbolické matematické výpočty na počítači. Nejznámějšími zástupci této skupiny jsou komerční programy Mathematica a Maple. Autor zmiňovaného článku mluví o méně známém, přitom však velmi kvalitním programu *MuPAD Light 2.0*¹, který je možno získat *zdarma* pro vzdělávací účely. Autor tvrdí, že program *MuPAD* lze využít při výuce na střední škole...

Dále však nechme mluvit jednoho z autorů. »V době svých studií na Matematicko-fyzikální fakultě jsem se poprvé seznámil s počítačovým programem *Mathematica*. Nadchly mě možnosti, které přináší pro zjednodušení a zkvalitnění práce v oblasti numerických i symbolických výpočtů. Vzpomínám, jak jsem jej s oblibou testoval. Počítal jsem například číslo π s přesností na několik stovek tisíc desetinných míst, řešil nejrůznější rovnice, hledal primitivní funkce a kreslil grafy často i velmi bizarních funkcí. Bavilo mě hledat nedokonalosti programu a shromažďovat příklady, při nichž počítačové řešení

Klíčová slova: *MuPAD*, systémy pro počítačovou algebru.

Poděkování: Tato práce vznikla za podpory grantu MSM 0021620839.

¹Vývoj v oblasti software je velmi rychlý, takže v době, kdy tento příspěvek jde do tisku, již jeho čtenář na Internetu místo verze 2.0 nalezne verzi 2.5.2. Domníváme se však, že všechny jeho základní rysy zůstávají tytéž, takže jsme zařadili pouze krátký odstavec informující o novinkách a změnách v *MuPADu* verze 2.5. oproti textu připraveného na základě našich zkušeností s používáním verze starší.

selhávalo. Později jsem pracoval i s dalšími matematickými systémy, z nichž mě nejvíce zaujal Maple a Matlab.

O to více mě mrzelo, že jsem nemohl žádný z těchto skvělých programů využívat ve své učitelské praxi. Učím na gymnáziu matematiku a je mi jasné, že běžná střední škola si nemůže zakoupit program, jehož multilicence pro jednu počítačovou učebnu stojí řádově několik set tisíc korun. Pokoušel jsem se trochu improvizovat a pátrat po levnějších alternativách. Z dob svých studií jsem znal ještě český program *Famulus*, se kterým se mi také velmi dobře pracovalo. Není sice vybaven symbolickými výpočty, nabízí však rychlé numerické algoritmy a snadno programovatelný grafický výstup. Vždy se mi líbil i jeho jednoduchý programovací jazyk se syntaxí podobnou Pascalu. Ani cena programu není vysoká, nehledě na to, že velké množství středních škol ho již vlastní. Vše má však jeden háček. Program je k dispozici pouze ve verzi pro operační systém DOS, která je navíc již několik let stará. Další vývoj programu byl zastaven a bohužel se s největší pravděpodobností nedočkáme verzí pro operační systémy typu MS Windows či Linux².

Na Internetu se mi ještě podařilo nalézt několik volně šiřitelných programů, například pro zobrazování třírozměrných grafů. Ve srovnání se špičkovými komerčními CAS systémy byly však jejich možnosti žalostné.

Pak se na mne přeci jen usmálo štěstí. Na Letní škole historie matematiky jsem si zakoupil CD ROM [1], na němž je zpracován kurz diferenciálního počtu funkcí více proměnných s využitím programu Maple. Jedna z kapitol textu se zabývá problematikou využití počítačů při výuce matematiky. Je doplněna poměrně obsáhlým přehledem matematického softwaru vhodného pro výuku. V něm jsem objevil odkaz na mě do té doby neznámý CAS systém – *MuPAD*. Z krátkého popisu jsem se dozvěděl, že systém je vyvíjen na univerzitě v Paderbornu v Německu. Nejvíce mě však zaujala věta: „*MuPAD* je nabízen po vyplnění licenčního ujednání *zdarma*.“ Neváhal jsem a navštívil domovskou stránku celého projektu na internetové adrese:

<http://www.mupad.de/>

Tam jsem našel několik důležitých informací. Program *MuPAD* existuje v několika verzích pro různé operační systémy. Z neznámějších uveďme MS Windows, Linux, Unix a Apple Macintosh OS. *Některé verze jsou placené,*

² *Linux* je volně šiřitelný operační systém. Od dob svého vzniku v roce 1991, kdy byl „hračkou“ skupiny počítačových nadšenců, urazil obrovský kus cesty vpřed. Dnes je plnohodnotným operačním systémem, vhodným i pro běžné uživatele. Na Internetu je možno zdarma získat nejen samotný Linux, ale i velké množství pro něj určených kvalitních aplikací. Jinou možností je zakoupení některé distribuce Linuxu na CD za symbolickou cenu 100 – 200 Kč.

*jiné jsou nabízeny po registraci zdarma*³. Například pro operační systém MS Windows jsou k dispozici dvě verze. Odlehčená verze *MuPAD Light*, která je zdarma, a ostrá verze *MuPAD Pro*, za kterou musíme v každém případě zaplatit. Pro Linux je k dispozici verze jediná, která je zdarma. Navíc si můžeme vybrat mezi anglickou a německou verzí programu. Můj zájem se soustředil na verze pro MS Windows, neboť tento operační systém používáme na naší škole nejvíce. Začal jsem tedy stahovat nejnovější anglickou verzi *Mupad Light 2.0 for Windows* a oddal se příjemnému pocitu očekávání něčeho nového a neznámého... «

2. MuPAD Light se představuje

Celá distribuce je zkomprimována v jediném spustitelném souboru o velikosti přibližně 10 MB. Instalace probíhá způsobem běžným ve Windows. Hardwarové nároky programu jsou poměrně nízké. Za minimální konfiguraci je možno považovat systém s procesorem 486/66 Mhz s 16 MB RAM a 25 MB volného prostoru na disku. Po zdárném ukončení instalace můžeme program ihned spustit. Přivítá nás hlavní okno se zobrazenou výzvou k zaregistrování programu⁴.

V horní části okna je k dispozici běžné roletové menu a tlačítková lišta s ikonami určenými k ovládání a změnám nastavení programu. Veškerý dialog s programem probíhá přes příkazovou řádku, což je v programech podobného typu běžné. Když si navíc zapamatujeme několik klávesových zkratk, potom při ovládání programu prakticky nebudeme potřebovat nabídkové menu ani tlačítkovou lištu. To přináší velkou výhodu zejména v rychlosti komunikace s programem.

Podívejme se, jak vypadá práce s programem v praxi. Představme si, že nás bude zajímat kolik je 2^{1234} . Potom stačí zapsat do příkazové řádky, která vždy začíná znaky >>:

2^{1234}

³V dalším textu ještě několikrát použiji formulaci: „Určitá verze *MuPADu* je k dispozici zdarma.“ Vždy tím myslím toto: Příslušnou verzi *MuPADu* můžeme po vyplnění licenčního ujednání využívat zdarma ve vzdělávacích institucích pro nekomerční účely. Další podrobnosti naleznete přímo v licenci, která je součástí každé distribuce *MuPADu* a přečíst si ji můžete i na internetové adrese <http://www.sciface.com/personal.shtml>

⁴Již zmíněná registrace je zdarma a probíhá formou vyplnění licenčního ujednání na domovských stránkách *MuPADu*. V neregistrované verzi je omezeno množství programem využívané paměti a náročnější výpočty tak často končí hláškou: „Výpočet přerušen z důvodu nedostatku paměti.“

Stiskem klávesy **Enter** odešleme zapsaný výraz ke zpracování. Okamžitě obdržíme přesný výsledek⁵:

```
29581122460809862906004469571610359078633968713537299223\
95562070506573507962389242610538372483780501864436477590\
70955993120820899330381760937027212482840944941362110665\
44377518349572681192920386118201521832389207735598339319\
12089288676526559936024879031137085494026686245211006117\
94270340232766099317098048887493809023127398253860618772\
619035009883272941129544640111837184
```

Dříve zapsaný text již není možno znovu editovat. Tato možnost je nabízena pouze v placené verzi programu *MuPAD* Pro. Pokud se chceme k zapsaným příkazům vrátit a například změnit některé parametry, jediným řešením je označit požadovaný text a zkopírovat jej přes schránku systému Windows, např. pomocí klávesové zkratky **CTRL + c**. Potom se přemístíme například pomocí klávesové zkratky **CTRL + End** na konec vstupní oblasti a zde stiskem **CTRL + v** vložíme text.

Přejdeme k další úloze. Máme určit druhou odmocninu z čísla 72. Žádný problém, jednoduše zadáme:

```
sqrt(72)
```

a získáme:

$$6 \sqrt{2}$$

Pokud vás výsledek překvapil, pak vezte, že *MuPAD* provádí všechny výpočty symbolicky a proto s absolutní přesností. Místo přibližného výsledku jsme obdrželi zjednodušený číselný výraz $6 \cdot 2^{\frac{1}{2}} = 6\sqrt{2}$. Chceme-li znát jeho vyjádření desetinným číslem, použijeme funkci `float()`:

```
float(%)
```

a obdržíme výsledek s přesností na 10 platných číslic:

```
8.485281374
```

Dodejme jen, že znak `%` má v jazyce *MuPADu* význam posledního výsledku.

Nestačí vám přesnost výsledku na deset platných číslic? Potom si ji nastavte třeba na sto platných číslic změnou hodnoty systémové proměnné **DIGITS** (pozor, musí být zapsáno velkými písmeny):

⁵Zpětná lomítka ve výpisu výsledku znamenají, že další cifry čísla pokračují na následujícím řádku.

```
DIGITS:= 100
```

Potom znovu zadáme:

```
float(sqrt(72))
```

a získáme výsledek s přesností na 100 platných číslic:

```
8.485281374238570292810132345258188471418031252261688439\
060078427944394870772642233102325205965849436
```

Zajímá vás například hodnota čísla π s přesností na 100 platných číslic? Pokud již máme nastaveno `DIGITS:=100`, zadáme pouze:

```
float(PI)
```

a za okamžik vidíme výsledek:

```
3.141592653589793238462643383279502884197169399375105820\
974944592307816406286208998628034825342117068
```

Poznamenejme, že přesnost výpočtů v *MuPADu* je omezena pouze paměťovými možnostmi počítače a dobou trvání výpočtu. Můžete si sami vyzkoušet, na kolik desetinných míst se vám podaří na vašem počítači číslo π určit.

- *MuPAD* obsahuje nepřeberné množství funkcí z mnoha oblastí matematiky. Zabýváte se třeba teorií čísel? Pak můžete využít například funkci, která zjišťuje, zda je dané číslo prvočíslem:

```
isprime(2398232343243249984321312321)
```

Výsledkem je pravdivostní hodnota:

FALSE

neboť zadané číslo prvočíslem není. Jestliž nás zajímá rozklad tohoto čísla na prvočinitele, stačí zadat:

```
factor(2398232343243249984321312321)
```

a obdržíme čtyři prvočinitele:

3 733 15117701 72140687161773179

Dobře, umíme rozložit přirozené číslo na prvočinitele. Co se však stane, pokud místo číselné hodnoty uvedeme v argumentu funkce `factor` například výraz $x^2 + 2x + 1$? Kdo by očekával, že program ohlásí chybu, bude překvapen. Po zadání:

```
factor(x^2 + 2*x + 1)
```

získáme výraz:

$$(x + 1)^2$$

který je rozkladem původního výrazu na součin. Tím jsme se dostali do oblasti symbolických výpočtů, ve které je *MuPAD* velmi silný.

Chcete nalézt například primitivní funkci k funkci $4x^3 - 5x + \sin(2x)$? Stačí zadat:

```
int(4*x^3 - 5*x + sin(2*x), x)
```

abychom dostali:

$$x^4 - \frac{5x^2}{2} - \frac{\cos(2x)}{2}$$

Užitečným nástrojem je i funkce `solve`. Umožňuje řešit celou řadu rovnic, včetně rovnic s parametrem a jejich soustav. Zvídavého čtenáře opět odkazují na článek [6], kde je tato funkce použita při řešení několika velmi hezkých středoškolských úloh.

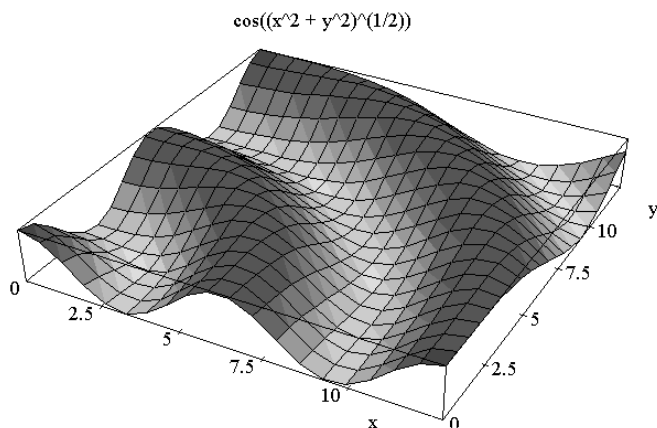
3. Grafika v MuPADu

Velmi efektní ukázkou možností *MuPADu* je také vykreslování grafických objektů. K dispozici máme procedury pro vykreslování 2D a 3D grafů funkcí, ale i obecnější procedury umožňující vizualizaci dat a nejrůznějších geometrických objektů. Samozřejmostí jsou i velké možnosti v nastavení popisu os a definování dalších vlastností grafických objektů, jako jsou například barvy, druhy čar apod.

Ukažme si, jak bude vypadat například část třírozměrného grafu funkce $f(x, y) = \cos(\sqrt{x^2 + y^2})$. Jestliže zadáme:

```
plotfunc3d(cos(sqrt(x^2 + y^2)), x = 0..4*PI,
y = 0..4*PI)
```


MuPAD otevře ve zvláštním okně prohlížeč grafů *VCam Light* se zobrazeným grafem (obrázek 1). Tlačítka v horní části okna můžeme graf libovolně otáčet, oddalovat i přibližovat a také měnit perspektivu pohledu.



Obrázek 1. Graf funkce $f(x, y) = \cos\left(\sqrt{x^2 + y^2}\right)$

Ukažme si, jak si *MuPAD* poradí se zobrazením parametricky zadaných křivek a ploch. Do jediného obrázku vykreslíme kulovou sféru spolu s jejími pravoúhlými průřezy do dvou navzájem kolmých promítacích rovin:

```
plot3d(Axes = None, Scaling = Constrained,
      // nezobrazovat osy, stejné měřítko na všech osách
      [Mode = Surface,
       [cos(u)*sin(v), sin(u)*sin(v), cos(v)],
       u~ = [0, 2*PI], v~ = [0, PI],
       // parametrická rovnice kulové sféry
       Grid = [15, 15], Smoothness = [2, 2],
       Style = [ColorPatches, AndMesh]
      ],
      [Mode = Surface,
       [u, v, -2],
       u~ = [-2, 2], v~ = [-2, 2],
       // parametrická rovnice půdorysné promítací roviny
       Grid = [15, 15], Smoothness = [2, 2],
       Style = [ColorPatches, AndMesh]
      ],
```

```

[Mode = Surface,
  [-2, u, v],
  u~ = [-2, 2], v~ = [-2, 2],
  // parametrická rovnice nárysne promítací roviny
  Grid = [15, 15], Smoothness = [2, 2],
  Style = [ColorPatches, AndMesh]
],

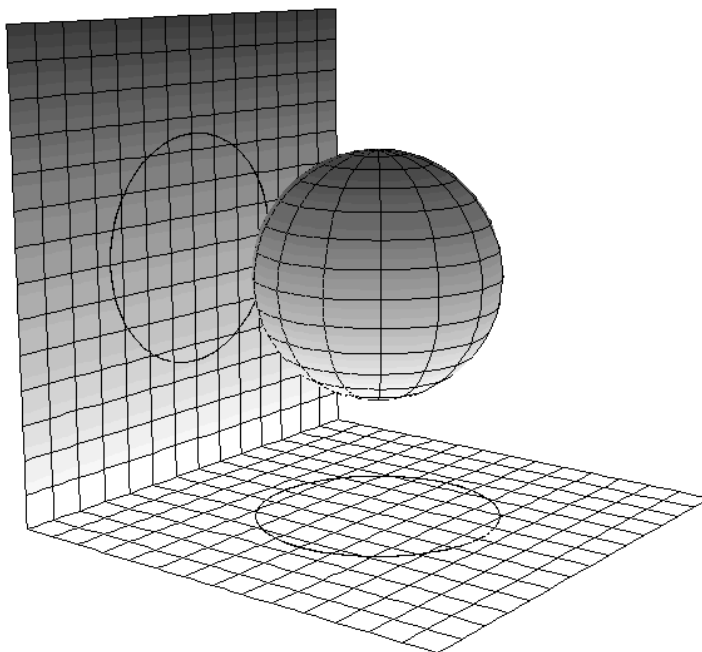
[Mode = Curve,
  [cos(u), sin(u), -2],
  u~ = [-2*PI, 2*PI],
  // parametrická rovnice kružnice v~půdorysně
  Grid = [100], Smoothness = [2],
  Color = [Flat, RGB::Black]
],

[Mode = Curve,
  [-2, cos(u), sin(u)],
  u~ = [-2*PI, 2*PI],
  // parametrická rovnice kružnice v~nárysne
  Grid = [100], Smoothness = [2],
  Color = [Flat, RGB::Black]
]
)

```

Výsledek vidíme na obrázku 2. Při popisu grafických objektů jsme kromě rovnic použili ještě speciální parametry určující způsob vykreslení. Parametrem **Grid** nastavujeme počet bodů, mezi kterými je prováděna interpolace křivek a ploch. Jestliže nastavíme vyšší počet bodů, získáme hladší křivky a hustší mřížku znázorňující prostorové plochy, avšak vykreslení objektů bude trvat déle. Pomocí parametru **Smoothness** můžeme bez zahuštění mřížky přidat další interpolační body, čímž ještě více vyhladíme křivky a plochy. Parametrem **Style** se nastavuje způsob zobrazení plochy. Volba **AndMesh** zajistí vykreslení kompletní povrchové mřížky a volba **ColorPatches** způsobí vybarvení jednotlivých povrchových plošek.

V dokumentaci nalezneme mnoho dalších parametrů sloužících k popisu způsobu zobrazení grafických objektů i celé scény. Například můžeme volit tzv. **CameraPoint**, což je bod, ze kterého celou scénu sledujeme, dále parametry perspektivního zobrazení, barvy jednotlivých křivek a ploch, nebo třeba způsob zobrazení soustavy souřadnic. Vše je možno doplnit popisky s volitelným fontem.



Obrázek 2. Zobrazení parametrických 3D křivek a ploch.

Vytvořenou scénu můžeme uložit do souboru ve speciálním formátu používaném prohlížečem. Bohužel verze VCam Light neumožňuje uložení obrázku v běžných grafických formátech *PostScript*, *JPEG*, *GIF* nebo *BMP*. Tato možnost je nabízena pouze v placené verzi programu *MuPAD Pro*. Je však samozřejmě možné získat bitmapové obrázky grafů běžným snímáním obrazovky ve Windows pomocí klávesy **Print Screen** a poté je zkonvertovat v některém vhodném grafickém editoru do požadovaného formátu⁶ Program je možné získat na adrese <http://www.xnview.com/>.

4. Vyšetřování průběhu funkce

Pojďme se nyní podívat na to, jak nám může být *MuPAD* nápomocen při řešení důležité úlohy ze středoškolské matematiky – vyšetřování průběhu funkce. Máme například zadánu funkci

⁶Podobným způsobem autor získal obrázky pro článek, který právě čtete. Snímání obrazovky bylo provedeno přímo z výborného freewarového grafického prohlížeče *XnView*, který zároveň umožňuje konverzi souboru do téměř libovolného grafického formátu.

$$f(x) = \frac{1}{x^2} - x.$$

Abychom nemuseli výraz, jímž je funkce definována, stále opisovat, označíme si ji identifikátorem `f`:

```
f := x -> 1/x^2 - x
```

Vyšetřování průběhu funkce zahájíme určením definičního oboru. Pomocí funkce `discont` nalezneme body, ve kterých není funkce definována:

```
discont(f(x), x)
```

```
{0}
```

Jistě nás bude zajímat průběh funkce v okolí bodu 0. Spočteme tedy limity pro $x \rightarrow 0_+$ a $x \rightarrow 0_-$:

```
limit(f(x), x = 0, Right), limit(f(x), x = 0, Left)
```

```
infinity, infinity
```

Podobně zjistíme, že limity pro $x \rightarrow \mp\infty$ jsou rovny $\pm\infty$:

```
limit(f(x), x = -infinity), limit(f(x), x = infinity)
```

```
infinity, -infinity
```

Průsečíky grafu funkce s osou x nalezneme řešením rovnice $f(x) = 0$:

```
solve(f(x) = 0, x)
```

```
{1, - 1/2 I~3^1/2 - 1/2, 1/2 I~3^1/2 - 1/2}
```

Jediným reálným kořenem rovnice je číslo 1. *MuPAD* našel ještě dva komplexní kořeny $-\frac{1}{2} - \frac{\sqrt{3}}{2}i$ a $-\frac{1}{2} + \frac{\sqrt{3}}{2}i$.

Nyní se budeme zabývat hledáním extrémů funkce. Nejprve vyjádříme první derivaci funkce:

```
f'(x)
```

$$-\frac{2}{3}x - 1$$

a ihned řešíme rovnici $f'(x) = 0$:

```
solve(f'(x) = 0, x, Domain = Dom::Real)
```

Klíčové slovo **Domain** označuje množinu, ve které chceme rovnici řešit. Protože nás zajímají pouze reálné kořeny, použijeme **Dom::Real**. Výsledkem budeme nejspíše zaskočení:

$$\{(-2)^{1/3} (1/2 - i\sqrt{3}/2) - 1/2\}$$

Očekávali jsme reálné číslo, obdrželi jsme však podivný výraz obsahující součin čísel $\sqrt[3]{-2}$ a $-\frac{1}{2} + \frac{\sqrt{3}}{2}i$. Podobná překvapení nás při práci s programy typu CAS čekají poměrně často. Musíme si uvědomit, že výraz $\sqrt[3]{-2}$ je matematickým zápisem úlohy řešit kubickou rovnici

$$x^3 + 2 = 0$$

v oboru komplexních čísel. Tato úloha má však tři řešení. Z důvodu jednoznačnosti zápisu v *MuPADu* je symbolem $\sqrt[3]{-2}$ zastoupen právě ten kořen rovnice, jenž má nejmenší argument v goniometrickém zápise. V našem případě je to kořen $\sqrt[3]{2}(\cos \frac{1}{3}\pi + i \sin \frac{1}{3}\pi)$. Jestliže jej s použitím Moivreovy věty vynásobíme číslem $-\frac{1}{2} + \frac{\sqrt{3}}{2}i = \cos \frac{2}{3}\pi + i \sin \frac{2}{3}\pi$ obdržíme výsledek $-\sqrt[3]{2}$.

A jak si usnadnit předchozí výpočty s *MuPADem*? Jestliže máme o výsledku jistou představu, můžeme *MuPADu* poradit, jaké úpravy má s výrazem provést, aby jej získal v jednodušším tvaru. V našem případě se nabízí použít funkci **rectform**, která nalezne reálnou a imaginární část komplexního výrazu z , tj. zapíše jej ve tvaru $z = \Re(z) + \Im(z)i$:

```
rectform(%)
```

$$\{-2^{1/3}\}$$

Vidíme, že řešením je skutečně reálné číslo $-\sqrt[3]{2}$, neboť imaginární část komplexního výrazu je nulová. Nalezli jsme lokální extrém funkce f . Hodnota druhé derivace funkce v bodě $-\sqrt[3]{2}$:

$$f''(-2^{1/3})$$

$$\frac{2/3}{3 \cdot 2} = \frac{1}{9}$$

je kladná a proto funkce v tomto bodě nabývá minima.

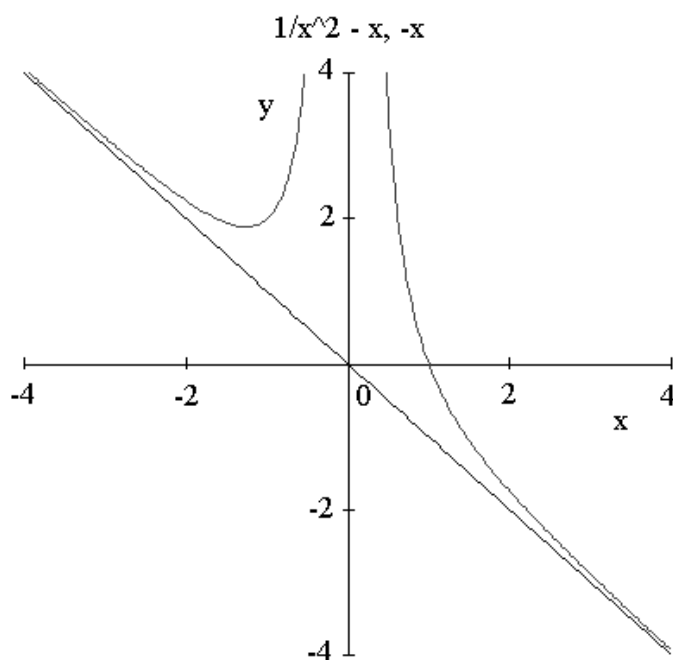
Zabývejme se ještě asymptotickým chováním funkce f . Pro $x \rightarrow \mp\infty$ se hodnota výrazu $1/x^2$ blíží k nule a je tedy zřejmé, že hodnoty funkce $f(x) = 1/x^2 - x$ se asymptoticky blíží k hodnotám funkce $g(x) = -x$:

$$g := x \rightarrow -x$$

Naše výsledky můžeme snadno ověřit vykreslením grafu funkce f spolu s asymptotou g :

```
plotfunc2d(f(x), g(x), x = -5..5, y = -5..5)
```

Zobrazený graf vidíme na obrázku 3.



Obrázek 3. Graf funkce $f(x) = \frac{1}{x^2} - x$.

5. Programování

Množství předdefinovaných funkcí jazyka *MuPADu* je obrovské. Přesto dříve či později narazíme na problém, který nelze řešit žádnou z vestavěných funkcí. Je běžné, že matematické systémy nabízejí pro tyto účely vhodný programovací jazyk. Jeho prostřednictvím je možno celý systém neustále zdokonalovat a rozšiřovat o tzv. *knihovny funkcí*, řešící problémy z různých oblastí matematiky, fyziky a dalších věd. V základní distribuci MuPAD Light je obsaženo 37 knihoven, mimo jiné z těchto oblastí: lineární algebra, diskretní matematika, kombinatorika, teorie čísel, statistika, lineární optimalizace, integrace, diferenciální rovnice, numerická matematika atd. Další knihovny můžeme získat od jiných uživatelů *MuPADu* například v internetových diskusních skupinách.

Z hlediska rozvoje našeho intelektu je však cennější chybějící funkci naprogramovat. *MuPAD* nám nabízí jazyk, jehož syntaxe je velmi podobná Pascalu. Domnívám se, že rozhodnutí tvůrců *MuPADu* bylo v tomto případě velmi šťastné. Pascal je v současnosti nejčastěji vyučovaným jazykem na středních i vysokých školách a ovládá jej nejvíce studentů i učitelů. Ti budou jistě potěšeni, že nemusí ztrácet čas učením se zcela novým příkazům a programovým konstrukcím.

Z jazyka *MuPADu* jsou oproti Pascalu vypuštěny některé prvky, které v matematickém systému nenacházejí využití (například typ *záznam*). Naproti tomu bylo přidáno množství užitečných datových struktur vhodných pro práci s matematickými objekty. Jmenujme například *seznamy*, *tabulky*, *vektory*, *matice*, či *polynomy*.

Ukažme si, jak můžeme naprogramovat funkci počítající faktoriál daného přirozeného čísla. Při zápisu zdrojového textu procedury přímo na vstup *MuPADu* není možné ukončovat řádky stiskem klávesy **Enter**, neboť by došlo k jejich okamžitému vyhodnocení. Pro přechod na další řádek je nutno použít kombinaci kláves **Shift + Enter** (tzv. měkký konec řádku). Dávku příkazů oddělených měkkými konci řádků najednou spustíme jediným stiskem klávesy **Enter**:

```
faktorial := proc(n : Type::NonNegInt)
begin
  if n = 0
  then return(1)
  else return(n*faktorial(n - 1))
  end_if
end_proc
```

Na prvním řádku přiřadíme funkci identifikátor `faktorial`, pomocí něhož bude později volána. Trochu matoucí může být použití klíčového slova `proc`. V *MuPADu* jej používáme nejen při definici procedur, ale i funkcí. Pro předání návratové hodnoty používáme funkci `return`. Všimněte si, že v jejím argumentu můžeme uvést i rekurentní volání právě definované funkce.

S takto definovanou funkcí pracujeme stejně jako s jinými funkcemi *MuPADu*. Chceme-li určit faktoriál čísla 100, stačí zadat:

```
faktorial(100)
```

```
93326215443944152681699238856266700490715968264381621468\
59296389521759999322991560894146397615651828625369792082\
722375825118521091686400000000000000000000000000000000
```

Pro první kroky programátora doporučujeme zejména výborný *MuPAD* Tutorial, který je součástí distribuce programu. S jeho pomocí se velmi rychle naučíte pracovat s programovými konstrukcemi a datovými strukturami⁷.

Určitou nevýhodou volně šířitelné verze Light je velmi omezené vývojové prostředí v němž chybí editor a debugger. Zdrojové texty procedur můžeme sice zapisovat přímo do příkazové řádky, bohužel však s minimálními možnostmi editace. Zejména v případě delších algoritmů je práce v příkazové řádce značně nepohodlná. Výhodnější je využít libovolný externí textový editor. Vybrat si můžeme třeba *Poznámkový blok* systému Windows, nebo jakýkoliv svůj oblíbený editor. Zapišeme-li v editoru zdrojový text nějaké funkce, máme několik možností jak jej přenést do systému *MuPAD*. Můžeme použít buď schránku systému Windows, nebo příkaz `read` k načtení souboru a definovanou funkci ihned vyzkoušet v interaktivním režimu. Jestliže budeme chtít funkci používat častěji, vyplatí se zařadit ji do některé z uživatelských knihoven, čímž ji budeme mít neustále k dispozici.

Pokud vám vadí nepřítomnost debuggeru, máme pro vás jeden tip. Ve volně šířitelné Linuxové verzi *MuPADu* je debugger k dispozici a tak klidně můžete ladit programy v Linuxu a hotové je prezentovat třeba v MS Windows.

Uvedme ještě jednu pro programátory nesmírně cennou informaci. Veškeré procedury obsažené v knihovnách *MuPADu* jsou k dispozici ve formě zdrojových textů v adresáři `lib`. Máme tak možnost zjistit, jakým způsobem jsou jednotlivé algoritmy implementovány a poučit se z nich při tvorbě vlastních knihoven. Je třeba si uvědomit, že není zdaleka běžné nabízet uživatelům volný přístup ke zdrojovým textům knihoven matematických programů. Naopak, v komerčních programech jsou knihovny obvykle zkompileovány. Autoři

⁷Tutorial je možno získat v anglické nebo německé verzi.

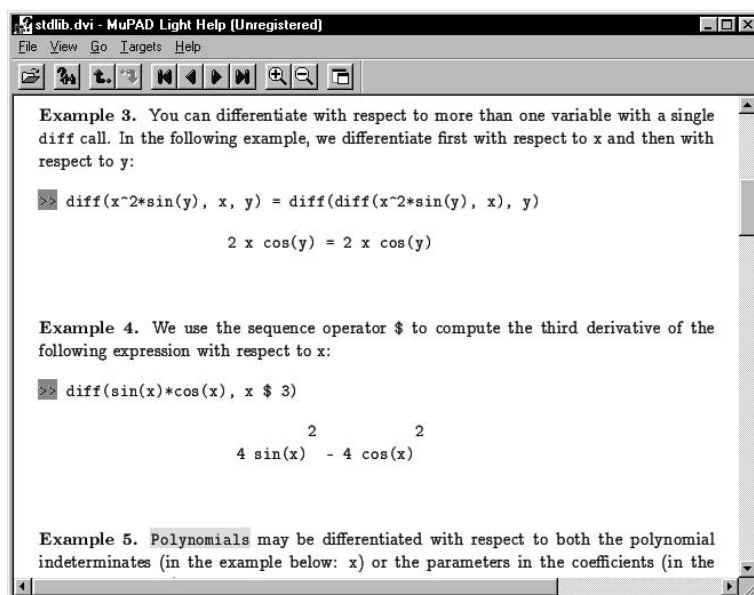
programů si tak chrání své duševní vlastnictví. Hlubavějším uživatelům však obvykle vadí, když jsou nuceni používat knihovny jako „černé skříňky“, které sice *něco* počítají, ale není možno zjistit *jak* to počítají. O to více je třeba v tomto směru ocenit přístup autorů *MuPADu*.

6. OnLine Help System

Na závěr jsme si nechali jednu z největších předností programu *MuPAD*, kterou je interaktivní systém nápovědy. Vyvolat jej můžeme přímo z příkazové řádky. Například nápovědu k funkci `diff` získáme takto:

? diff

Nejprve se otevře tzv. *Help Browser*, což je běžný rejstřík známý z nápovědy v systému Windows. Zde upřesníme název hledaného hesla a po potvrzení se otevře hypertextový prohlížeč s požadovanou manuálovou stránkou (obrázek 4).



Obrázek 4. Interaktivní nápověda systému MuPad.

Aktivní hesla jsou v dokumentu podbarvena žlutě. Jestliže na ně klikneme myší, přeneseme se na další stránky s upřesňující nápovědou. Kromě

toho máme k dispozici několik tlačítek pro pohyb na stránkách. Za nejužitečnější považují tlačítka pro přesun vpřed a vzad v historii dosud zobrazených stránek.

Popis každé funkce je v manuálu doplněn spoustou ukázek použití. Vedle každé z nich je umístěn zeleně podbarvený znak `>>`. Jestliže na něj klikneme, zdrojový text ukázky se přenese přímo do příkazové řádky *MuPADu*. Takto můžeme při čtení manuálu velmi rychle a pohodlně spouštět veškeré příklady a ihned si prohlížet výsledky. Výhody interaktivní práce s nápovědou vyniknou zejména při studiu tutorialu, který je dodáván spolu s programem. Po několika hodinách práce s ním jsem došel k názoru, že si takto představuji vzorovou interaktivní učebnici.

Je vcelku pochopitelné, že uživatel zatouží vytvořit si vlastní hypertextovou učebnici spolupracující s *MuPADem*, která by se dala využít ve školní výuce. Po chvíli pátrání v dokumentaci zjistí, že program zobrazující manuálové stránky je „obyčejný“ prohlížeč `dvi`⁸ souborů obohacený o hypertextové funkce a propojení s *MuPADem*.

Někteří čtenáři systém $\text{\TeX}$ možná neznají a proto nám dovoluňte uvést o něm několik základních informací. Jedná se o volně šířitelný systém, který se po celém světě využívá při sazbě vědeckých časopisů a knih. Například všechny příspěvky pro váš kurz byly vysázeny v $\text{\TeXu}$. Vědci jej mají ve velké oblibě, neboť umožňuje precizní sazbu i velmi složitých matematických vzorců a rovnic. V $\text{\TeXu}$ lze však velmi dobře sázet i chemické vzorce, elektronická schémata nebo třeba notové záznamy. Autorem $\text{\TeXu}$ je profesor Donald E. Knuth ze Stanfordské univerzity, jeden z nejznámějších vědců v oboru výpočetní techniky. Systém existuje již téměř dvacet let, jeho geniální koncepce však dodnes zaznamenala minimálních změn. K systému je v současné době k dispozici velké množství nadstaveb, z nichž nejznámější je soubor maker $\text{\LaTeX}$. S jeho pomocí zvládnou i běžní uživatelé, kteří nejsou seznámeni s typografickými zásadami, vytvořit dokumenty dokonalého vzhledu. Chcete-li se dozvědět o $\text{\TeXu}$ a všem co s ním souvisí více, doporučujeme například knihu [5]. Pro seznámení s $\text{\LaTeXem}$ je vhodná kniha [8]. Obrovské množství informací dále naléznete i na serveru Československého sdružení uživatelů $\text{\TeXu}$:

<http://www.cstug.cz/>

kde naleznete i různé distribuce $\text{\TeXu}$ spolu s množstvím maker a pomocných programů.

⁸`dvi` je výstupní formát systému $\text{\TeX}$. Jedná se o zkratku anglického *DeVice Independent*, což znamená výstupní formát nezávislý na zobrazovacím zařízení.

Vraťme se ale k systému *MuPAD*. Postup tvorby interaktivní učebnice je velice snadný. Vytvoříme běžný dokument v $\text{\LaTeX}$ u a pomocí balíku `mhHelp` jej opatříme hypertextovými odkazy a značkami. Změnou jediného parametru v hlavičce dokumentu potom můžeme zdrojový soubor přeložit $\text{\TeX}$ em buď na běžný `dvi` soubor, který je vhodný pro tisk, nebo na `dvi` soubor s hypertextovými vlastnostmi, který je určen pro interaktivní práci v *MuPAD*u. Pro zajímavost uvedme, že i tento článek byl vytvořen stejným způsobem. Je tedy zbytečné, aby jste při svých pokusech s *MuPAD*em opisovali zde uvedené ilustrační příklady. Raději si stáhněte interaktivní verzi článku v podobě `dvi` souboru z internetové adresy:

<http://www.cihak.com/mupad/mupad2004.dvi>

Zde naleznete i několik dalších ukázek použití *MuPAD*u ve výuce matematiky a také české fonty pro hypertextový prohlížeč *MuPAD*u potřebné pro správné zobrazení dokumentů v češtině.

Přímé propojení precizně vysázeného $\text{\TeX}$ ovského dokumentu s *MuPAD*em lze považovat za jednu z největších předností tohoto programu. Když k ní připočteme rychlé výpočetní jádro, výborný a snadno zvládnutelný programovací jazyk a v neposlední řadě také efektní grafický výstup, nezbyvá než konstatovat, že se jedná o špičkový produkt zejména pro potřeby výuky a vzdělávání. Učitelé tak dostávají *zdarma* do rukou účinný soubor nástrojů pro tvorbu velmi kvalitních interaktivních matematických učebnic nejrůznějšího druhu, od jednoduchých lekcí určených středoškolákům až po přípravu rozsáhlých učebních textů zahrnujících látku některého univerzitního kurzu.

Volně šířitelná verze programu má sice oproti verzi placené některá omezení jako například chybějící notebook a debugger, nebo omezené možnosti ukládání grafiky, ty však nikterak nesnižují jeho užitnou hodnotu v oblasti tvorby interaktivních učebnic.

Dalším neopomenutelným kladem programu je fakt, že je neustále vyvíjen a zdokonalován. Dokladem toho je nejen nová verze 2.5, která byla dokončena teprve před několika měsíci, ale i množství projektů zabývajících se jeho využitím v různých oblastech vědy a vzdělávání. Na domovské stránce *MuPAD*u můžeme nalézt odkazy na několik velmi zajímavých vzdělávacích projektů, které probíhají na vysokých i středních školách v Německu a ukazují možnosti počítači podporované výuky matematiky.

Závěrem si dovoluujeme nabídnout pomoc všem zájemcům o tvorbu a využití interaktivních výukových materiálů v *MuPAD*u. Neváhejte a napište o svých plánech na adresu:

mcihak@karlin.mff.cuni.cz

Na oplátku obratem obdržíte potřebná makra pro $\text{\LaTeX}$ včetně návodu na jejich použití v dokumentech. Pokud se budete chtít pochlubit s výsledky své tvorby ostatním uživatelům *MuPADu*, nabízíme možnost uveřejnění vašich prezentací na internetové adrese:

<http://www.cihak.com/mupad/>

kde by jeden z autorů chtěl postupně vytvářet archiv českých výukových materiálů vhodných pro použití v hodinách matematiky.

Zajímavé internetové adresy

- <http://www.mupad.de/> – domácí stránky projektu *MuPAD*. Zde je možné získat distribuce *MuPADu* pro různé operační systémy, velké množství dokumentace v angličtině nebo v němčině, informace o nových projektech a adresy internetových diskusních skupin zabývajících se *MuPADem*.
- <http://www.cihak.com/mupad/> – stránky jednoho z autorů tohoto textu. Ukázky interaktivních výukových materiálů pro program *MuPAD*. České fonty pro hypertextový prohlížeč textů a helpů *MuPADu*.
- <http://www.cstug.cz/> – stránky Československého sdružení uživatelů $\text{\TeX}$ u. Naleznete zde mnoho užitečných informací o typografickém systému $\text{\TeX}$, množství maker a pomocných programů, odkazy na literaturu a další zajímavá místa na internetu.
- <http://www.xnview.com/> – stránky grafického prohlížeče *XnView*.
- <http://www.mathworks.com/> – Internetové stránky producenta systému *Matlab*.
- <http://www.wolfram.com/> – Internetové stránky producenta systému *Mathematica*.

Literatura

- [1] Došlá Z., Plch R. a Sojka P. (1999). *CD ROM Matematická analýza s programem Maple: Diferenciální počet funkcí více proměnných*. Masarykova univerzita, Brno.
- [2] Dvořák L., Ledvinka T., Sobotka M. (1991). *Famulus 3.1 – příručka uživatele*, Praha.
- [3] Oevel W. a kol. (2000). *The MuPAD Tutorial*. Elektronická dokumentace, SciFace.
- [4] Oevel W., Gerhard J. (2000). *The MuPAD 2.0 Quick Reference*. Elektronická dokumentace, SciFace, 2000.
- [5] Olšák P. (1999). *Typografický systém T_EX*. Konvoj, Brno.
- [6] Lukáč S. (2002). *Programy typu CAS vo vyučování matematiky*, MFI, **12(1)**.
- [7] Plch R. (1997). *Internet pro učitele matematiky*. Prometheus, Praha.
- [8] Rybička J. (1995). *L^AT_EX pro začátečníky*. Prometheus, Praha.
- [9] Wolfram S. (1990). *Mathematica, A system for Doing Mathematics by Computer*, 2. vydání. Addison–Wesley, 1991.

Appendix A. Pracovní materiál pro přednášku

Během přednášky si použití *MuPADu* předvedeme na následujících příkladech, které si podrobně okomentujeme.

HELP

`info(ln)`

`?solve`

`info(stdlib)`

`info(solverlib)`

ZÁKLADNÍ FUNKCE A OPERÁTORY

Funkce	Popis funkce
<code>+, -, *, /, ^</code>	základní aritmetické operace
<code>abs</code>	absolutní hodnota
<code>ceil</code>	zaokrouhlení „nahoru“
<code>div</code>	celá část po operaci „modulo“
<code>fact, !</code>	faktoriál
<code>float</code>	aproximace reálným číslem
<code>floor</code>	zaokrouhlení „dolů“
<code>frac</code>	zlomková část
<code>ifactor, factor</code>	rozklad na prvočinitele
<code>isprime</code>	test prvočíselnosti
<code>mod</code>	zbytek po operaci „modulo“
<code>round</code>	zaokrouhlení
<code>sign</code>	znaménko
<code>sqrt</code>	odmocnina
<code>trunc</code>	celá část

Základní funkce a operátory pro čísla v *MuPADu*.

„PŘESNÉ“ POČÍTÁNÍ S ČÍSLY

```
1 + 5/2
(1 + (5/2 * 3))/(1/7 + 7/9)^2
sqrt(28)
2^100
2^10000
100!
isprime(123456789)
ifactor(123456789)
factor(123456789)
```

„NEPŘESNÉ“ POČÍTÁNÍ S ČÍSLY

```
float(sqrt(28))
sqrt(28.0)
sqrt(28.)
(1 + (5/2 * 3))/(1/7 + 7/9)^2
(1.0 + (5/2 * 3))/(1/7 + 7/9)^2
DIGITS; float(PI)
DIGITS; float(E)
cos(PI), ln(E)
DIGITS := 100 : float(PI); float(E); DIGITS := 10 :
```

```

2/3*sin(2), 0.6666666666*sin(2)

float(2/3 * sin(2)), float(0.6666666666 * sin(2))

sqrt(-1), I^2

(1 + 2*I)*(4 + I)

1/(sqrt(2) + I)

rectform(1/(sqrt(2) + I))

Re(1/(sqrt(2) + I)), Im(1/(sqrt(2) + I))

abs(1/(sqrt(2) + I)), conjugate(1/(sqrt(2) + I)),
rectform(conjugate(1/(sqrt(2) + I)))

```

SYMBOLICKÉ POČÍTÁNÍ

```

f := y^2 + 4*x + 6*x^2 + 4*x^3 + x^4

diff(f,x), diff(f,y)

diff(diff(diff(f, x), x), x), diff(f, x, x, x)

sin', sin'(x), D(sin), D(sin)(x)

int(f, x)

int(f, x = 0..1)

integral := int(1/(exp(x^2) + 1), x)

diff(integral, x)

int(1/(exp(x^2) + 1), x = 0..1)

float(%), last(2)

sum(i, i~= 1..n)

```



```
sum(i, i~= 1..n) : factor(%)
sum(i^2, i~= 1..n)
sum(1/i^2, i~= 1..infinity)
product(i^3, i= 1..n); rewrite(%, fact)
```

ZJEDNODUŠOVÁNÍ VÝRAZŮ

Funkce	Popis funkce
collect	koeficienty
combine	kombinace podvýrazů
expand	rozvoj (expanze)
factor	faktORIZACE
normal	nalezení společného jmenovatele a zkrácení „společných členů“
partfrac	partial fraction (rozklad na součet racionálních členů)
radsimp	zjednodušení racionálních výrazů
rectform	reálná a imaginární část komplexního výrazu
rewrite	zjednodušení výrazů pomocí identit
simplify	univerzální „zjednodušovač“

Základní funkce pro zjednodušování výrazů.

```
x^2 + a*x + x*sqrt(2) + sin(x) + a*sin(x)
collect(%, x)
collect(%, [x, sin(x)])
expand(exp(x + y)), expand(sin(x + y)),
  expand(tan(x + 3*PI/2))
factor(x^3 + 3*x^2 + 3*x + 1),
  factor(x^2/(x + y) - z^2/(x + y))
factor((exp(x)^2 - 1)/(sin(x)^2 - cos(x)^2))
```

```
f := ((x + 6)^2 - 17)/(x - 1)/(x + 1) + 1 :
f, factor(f), normal(f)

normal(x^2/(x + y) - y^2/(x + y))

f := x/(1 + x) - 2/(1 - x), g = normal(f),
partfrac(g, x)

simplify((exp(x) - 1)/(exp(x/2) + 1))

radsimp(sqrt(4 + 2*sqrt(3)))
```

ŘEŠENÍ (SOUSTAV) ROVNIC

```
solve(x^2 - 2*x + 2 = 0, x)

equation := {x + y = a, x - a*y = b} :
unknowns := {x,y} : solve(equation, unknowns);
```

POČÍTÁNÍ LIMIT

```
limit(sin(x)/x, x=0)

limit((1 + 1/n)^n, n = infinity)
```

GRAFIKA

Všechna volání funkcí pro dvourozměrné kreslení vypadají v zásadě následovně :

```
>> plot2d(sceneOptions1, sceneOptions2, ..., Object1,
           Object2, ...)
>> plot3d(sceneOptions1, sceneOptions2, ..., Object1,
           Object2, ...)
```

Mezi parametry „scény grafu“ mimo jiné patří parametry os (čára nebo šipka, jejich škálování, popis, font...), popis obrázku, barvy, mřížka, „hladkost“ vykreslení, velikost grafických objektů (markerů) apod.

Kreslení funkcí

```
plotfunc2d(sin(x^2), x = -2..5)

plotfunc2d(sin(x), cos(x), x=0..4*PI)

plotfunc2d(1/(1-x) + 1/(1+x), x=-2..2)

plotfunc3d(sin(x^2 + y^2), x = 0..PI, y = 0..PI)
```

Kreslení křivek

Parametrické křivky v 2D, například kruh, nemůžeme malovat stejně jednoduše jako funkce, nicméně i zde *MuPAD* nabízí snadné řešení, které umožňuje kreslit i funkce.

```
Object := [Mode = Curve, [x[u], y[u]], u~ = [a, b],
           objectOptions1, objectOptions2, ...]

SinGraf := [Mode = Curve, [u, sin(u)], u~ = [0, 4*PI],
           Grid = [100]] :

Circle1 := [Mode = Curve, [cos(u), sin(u)],
           u~ = [0, 2*PI]] :

sceneOptions1 := Background = [1,1,1], ForeGround =
           [0,0,0], Ticks = 0, Axes = Box :

sceneOptions2 := Background = [0,0,0], ForeGround =
           [1,1,1], Ticks = 0, Axes = Box :

plot2d(sceneOptions1, SinGraf)

plot2d(sceneOptions2, Circle1)

info(RGB)

export(RGB)

Circle2 := [Mode = Curve, [cos(u), sin(u)],
           u~ = [0, 2*PI], Grid = [10], Title = "Kružnice 1",
           TitlePosition = [1.7,2], Color = [Flat, RGB::Red]]:
```

```

Circle3 := [Mode = Curve, [cos(u)/2, sin(u)/2],
  u~ = [0, 2*PI], Grid = [100], Title = "Kružnice 2",
  TitlePosition = [6.3, 2.8], Color =
  [Flat, RGB::Blue]] :

sceneOptions3 := Title = "Dvě kružnice", FontSize =
  14, TitlePosition = [5, 0.7], Axes = Origin,
  Labeling = TRUE, Labels = ["x", "y"],
  BackGround = RGB::Gold, ForeGround = [1,1,1]

plot2d(sceneOptions3, Circle2, Circle3)

```

Kreslení povrchů

Parametrické plochy v 3D, například kouli, nemůžeme malovat stejně jednoduše jako 3D-funkce, nicméně i zde *MuPAD* nabízí snadné řešení, které umožňuje kreslit i 3D-funkce.

```

Object := [Mode = Surface, [x(u,v), y(u,v), z(u,v)],
  u~ = [a, b], v~ = [A, B], objectOptions1,
  objectOptions2, ...]

plot3d(BackGround = [1,1,1], ForeGround = [0,0,0],
  Axes = None, [Mode = Surface, [cos(u) * sin(v),
  sin(u) * sin(v), cos(v)], u~ = [0,2*PI], v~ = [0,PI],
  Grid = [20,20], Smoothness = [2,2], Color=[Height],
  Style = [ColorPatches, AndMesh]])

plotfunc3d(sin(x^2 + y^2), x = 0..2*PI, y = 0..2*PI)

plot3d(BackGround = [1,1,1], ForeGround = [0,0,0],
  Axes = None, Title = "Graf funkce sin(x^2 + y^2)",
  TitlePosition = Below, [Mode = Surface, [x, y,
  sin(x^2 + y^2)], x = [0, 2*PI], y = [0, 2*PI],
  Grid = [20, 20], Smoothness = [2, 2],
  Color = [Height], Style = [ColorPatches, AndMesh]])

plot3d(Axes = Box, Ticks = 0, Scaling = Constrained,
  Title = "Spiral", TitlePosition = Below,
  [Mode = Curve, [cos(12*u*PI)*sin(u*PI),
  sin(12*u*PI)*sin(u*PI), cos(u*PI)], u~ = [0,1],
  Grid = [50], Smoothness = [5]])

```

```
plot3d(Axes = Box, BackGround = RGB::White,
  ForeGround = RGB::Black, [Mode = Surface, [x, y,
  x^2 - y^2], x = [-1, 1], y = [-1, 1], Grid =
  [15,15], Smoothness = [3,3], Style = [ColorPatches,
  AndMesh]])
```

Kreslení seznamů

```
PlotPoints := [point(i, sin(i*6.28/50)) $ i = 0..50] :
plot2d(BackGround = [1,1,1], ForeGround = [0,0,0],
  Scaling = UnConstrained, PointWidth = 30,
  [Mode = List, PlotPoints])
```

Kreslení vrstevnic, rotačních ploch apod.

```
info(plot)

MojePlocha := [x, y, (x^2 + y^2)^3 - (x^2 - y^2)^2],
  x = [-1, 1], y = [-1,1]

plot3d(Axes = Box, BackGround = RGB::White, ForeGround
  = RGB::Black, [Mode = Surface, MojePlocha, Grid =
  [15,15], Smoothness = [3,3], Style = [ColorPatches,
  AndMesh]])

plot(plot::contour(MojePlocha, Contours =
  [-0.05,0,0.05], Grid = [10,10]))

plot(plot::contour(MojePlocha, Contours =
  [-0.05,0,0.05], Grid = [100,100]))

r := plot::xrotate(sin(x), x = 1..2*PI) : plot(r)

plot(plot::xrotate(cos(x), x = 1..4*PI))
```

VLASTNÍ FUNKCE

```
F := x -> x^2 : F(x), F(y), F(a + b), F'(x)
```

JEDEN ELEMENTÁRNÍ PŘÍKLAD

Prostudujme si chování funkce

$$\frac{(x-1)^2}{(x-2)} + a$$

```
f := x -> ((x - 1)^2 / (x - 2) + a):
```

```
singularities := discontinuity(f(x), x)
```

```
limit(f(x), x = 2, Left), limit(f(x), x = 2, Right)
```

```
roots := solve(f(x) = 0, x)
```

```
f'(x)
```

```
extrema := solve(f'(x) = 0, x)
```

```
f''(1), f''(3)
```

```
f(1), f(3)
```

```
limit(f(x), x = -infinity), limit(f(x), x = infinity)
```

```
series(f(x), x = infinity)
```

```
F := subs(f(x), a~ = -4): G := subs(f(x), a~ = 0):
```

```
H := subs(f(x), a~ = 4):
```

```
F, G, H
```

```
plotfunc2d(F, G, H, x = -1..4)
```

HRÁTKY S PRVOČÍSLY

Podobně jako jiné systémy pro symbolickou matematiku poskytuje *MuPAD* řadu elementárních funkcí z oblasti teorie čísel, například:

Funkce	Popis funkce
isprime(n)	test prvočíselnosti
ithprime(n)	vrátí n -té prvočíslo
nextprime(n)	vrátí nejmenší prvočíslo větší než n
ifactor(n)	rozklad čísla n na prvočinitele

Vybrané elementárních funkce z oblasti teorie čísel.

```
primes := select([$ 1..10000], isprime)
nops(primes)

primes1 := [ithprime(i) $ i = 1..1229]
nops(primes1)

primes3 := select([ithprime(i) $ i = 1..5000],
  x -> (x <= 10000))
bool(primes1=primes3)

primes4 := []: i~:= 2 :
repeat
  primes4 := primes4 . [i];
  i~:= nextprime(i + 1);
until i~> 10000 end_repeat

primes := select([$ 1..500], isprime) :
dist1 := [primes[i]-primes[i-1] $ i=2..nops(primes)]
```

Funkce `zip(a, b, f)` kombinuje seznamy

$$\mathbf{a} = [a_1, a_2, \dots] \quad \mathbf{a} \quad \mathbf{b} = [b_1, b_2, \dots]$$

po složkách, přičemž na jednotlivé dvojice aplikuje funkci f . Jinými slovy, funkce `zip(a, b, f)` vrací seznam

$$[f(a_1, b_1), f(a_2, b_2), \dots]$$

a má tolik prvků, kolik jich má kratší seznam.

```

b := primes : delete(b[1]) :
dist2 := zip(b, primes, (x, y) -> (y-x))

bool(dist1 = dist2)

nops(dist1), dist1

nops(dist2), dist2

```

KNIHOVNY FUNKCÍ A PROCEDUR

Většina matematických znalostí není obsažena v jádru *MuPADu*, nýbrž v jednotlivých knihovnách, například pro:

Název knihovny	Použití
combinat	kombinatorika
Dom	předdefinované datové struktury
fp	funkcionální programování
generate	generátor výstupů pro C, Fortran a T _E X
linalg	lineární algebra
misc	různé
module	správa modulů
network	teorie grafů
numeric	numerické algoritmy
numlib	teorie čísel
plot	rutiny pro kreslení
property	vlastnosti identifikátorů
RGB	nastavení barev
stdlib	ZÁKLADNÍ STANDARDNÍ KNIHOVNA
student	vybrané elementární algoritmy
Type	identifikátory datových typů

Vybrané knihovny *MuPADu*.

Knihovna pro numerickou matematiku

```
info(numlib)

?numlib

expose(sin)

expose(numlib::quadrature)

decimal(2/3)

numeric::solve(sin(x), x = 2..4)

quadrature(exp(x) + 1, x = 0..1)

numeric::quadrature(exp(x) + 1, x = 0..1)

quadrature(exp(x) + 1, x = 0..1)

export(numeric, quadrature) :
    quadrature(exp(x) + 1, x = 0..1)

delete(quadrature) : quadrature(exp(x) + 1, x = 0..1)
```

JAK SKONČIT?

```
quit
```

ÚLOHY PRO CVIČENÍ

Ověřte pomocí *MuPADu*, že platí:

1.

$$\begin{aligned} \lim_{x \rightarrow 0} \frac{\sin(x)}{x} &= 1, & \lim_{x \rightarrow 0} \frac{1 - \cos(x)}{x} &= 0, \\ \lim_{x \rightarrow 0+} \ln(x) &= -\infty, & \lim_{x \rightarrow 0} x^{\sin(x)} &= 1, \\ \lim_{x \rightarrow \infty} \left(1 + \frac{1}{x}\right)^x &= e, & \lim_{x \rightarrow \infty} \frac{\ln(x)}{e^x} &= 0, \\ \lim_{x \rightarrow 0} x^{\ln(x)} &= \infty, & \lim_{x \rightarrow 0} \left(1 + \frac{\pi}{x}\right)^x &= e^\pi, \\ \lim_{x \rightarrow 0-} \frac{2}{1 + e^{-1/x}} &= 0, & \lim_{x \rightarrow 0} ((\sin(x)))^{1/x} &. \end{aligned}$$

2. Prostudujte *help* pro funkci `diff` a spočtěte pomocí *MuPADu* pátou derivaci funkce $\sin(x^2)$.

3. Spočtěte přesně $\sqrt{27} - 2\sqrt{3}$ a $\cos(\pi/8)$. Dále určete numerickou aproximaci obou výrazů s přesností na pět platných číslic.

4. Rozviňte výraz $(x^2 + y)^5$.

5. Ověřte pomocí *MuPADu*, že platí:

$$\frac{x^2 - 1}{x + 1} = x - 1.$$

6. Prohlédněte si *help* pro funkci `sum` a ověřte pomocí *MuPADu*, že platí:

$$\sum_{i=1}^n (i^2 + i + 1) = \frac{n(n^2 + 3n + 5)}{3}$$

7. Spočtete

$$\sum_{i=0}^{\infty} \frac{2k-3}{(k+1)(k+2)(k+3)} \quad \text{a} \quad \sum_{i=2}^{\infty} \frac{k}{(k-1)^2(k+1)^2}.$$

8. Nakreslete graf funkce

$$f(x) = \frac{1}{\sin(x)}, \quad 1 \leq x \leq 10.$$

9. Spočtete pomocí *MuPADu*

$$\int x \, dx, \quad \int_1^4 x \, dx \quad \text{a} \quad \frac{d}{dx} \ln \ln x.$$

10. Vyřešte pomocí *MuPADu* rovnici

$$x^4 + px + 1 = 0$$

v x s jedním parametrem p .

11. Vyřešte pomocí *MuPADu* soustavu rovnic

$$x + y = 1, \quad x - y = 1.$$

12. Co dostanete, napíšete-li v *MuPADu*

$$\text{assume}(x > 2): \text{is}(x^2 > 4), \text{is}(x^3 < 0), \text{is}(x^4 > 17)?$$

Appendix B. Přehled základních funkcí MuPADu

Tento dodatek je podstatně rozšířeným překladem textu W. Oevela, *MuPAD – An Overview*, 1998.

1. Informace a nápověda

? : vyvolání nápovědy
 error : výstup chybového hlášení
 info() : vypis krátké informace (o objektech,
 : datových typech,...)
 setuserinfo : nastavení informační úrovně
 userinfo : výpis informací o proceduře
 warning : výstup varování

2. Systémové proměnné

Systémové proměnné jsou globální proměnné, jež definují vlastnosti prostředí *MuPADu* a řídí operace algoritmů v procedurách *MuPADu*. Přednastavené (default) hodnoty mohou být uživatelem změněny.

	default	
DIGITS	10	počet platných číslic pro zobrazení výsledků
HISTORY	20	počet výsledků dostupných pomocí příkazu <code>last</code>
LEVEL	100	hloubka výpočtu
LIBPATH		cesta k souborům <i>MuPADu</i> a ke knihovnám
MAXDEPTH	500	maximální hloubka rekurze
MAXLEVEL	100	maximální hloubka výpočtu
ORDER	6	počet členů rozvoje funkce v řadě
PRETTYPRINT	TRUE	kontrola tvaru textového výstupu
READPATH		cesta k souborům pro čtení pomocí funkce <code>read</code>
SEED	1	počáteční hodnota pro generátor náhodných čísel, viz funkce <code>random</code>
TEXTWIDTH	75	maximální počet znaků na jednom řádku výstupu
WRITEPATH		cesta k souborům pro zápis pomocí funkce <code>write</code>

3. Předdefinované datové typy (Domény)

V jádru *MuPADu* jsou předdefinovány následující datové typy (domény):

DOM_ARRAY	: pole
DOM_BOOL	: logické konstanty TRUE, FALSE a UNKNOWN
DOM_COMPLEX	: komplexní čísla
DOM_DOMAIN	: datové struktury
DOM_EXEC	: jednoduché funkce
DOM_EXPR	: výrazy
DOM_FAIL	: objekt FAIL
DOM_FLOAT	: reálná čísla s plovoucí desetinnou čárkou
DOM_FUNC_ENV	: prostředí funkcí
DOM_IDENT	: jména objektů jazyka
DOM_INT	: celá čísla
DOM_LIST	: seznamy
DOM_NIL	: objekt NIL
DOM_NULL	: prázdný objekt null()
DOM_POINT	: body (jako základní grafický prvek)
DOM_POLY	: polynomy
DOM_POLYGON	: mnohoúhelníky (jako základní grafický prvek)
DOM_PROC	: procedury
DOM_RAT	: racionální čísla
DOM_SET	: nastavení
DOM_STRING	: textové řetězce
DOM_TABLE	: tabulky
DOM_VAR	: lokální proměnné procedur

Více datových struktur a informace o jejich konstrukci jsou dostupné v knihovně Dom (informace lze získat pomocí příkazu `info(Dom)`).

4. Generování datových struktur

->	: definuje proceduru (DOM_PROC) ⁹
array	: generuje pole (DOM_ARRAY)
asympt	: asymptotický rozvoj (Series::gseries)
factor, ifactor	: rozklad na součin nerozložitelných prvků (Factored)

⁹V závorce je vždy uveden datový (doménový) typ generované datové struktury.

<code>funcenv</code>	: definuje prostředí funkcí (<code>DOM_FUNC_ENV</code>)
<code>matrix</code>	: definuje matici nebo vektor (<code>Dom::Matrix()</code>)
<code>new</code>	: generuje nový datový typ (doménu)
<code>newDomain</code>	: generuje novou doménu (<code>DOM_DOMAIN</code>)
<code>null</code>	: generuje „prázdný objekt“ (<code>DOM_NULL</code>)
<code>0</code>	: generuje řád zbytku v Taylorově rozvoji
<code>ode</code>	: reprezentuje obyčejnou diferenciální rovnici
<code>numlib::contfrac</code>	: počítá rozvoj v řetězový zlomek (generuje <code>numlib::contfrac</code>)
<code>piecewise</code>	: generuje podmíněně definovaný objekt
<code>point</code>	: generuje datovou strukturu <code>DOM_POINT</code> pro kreslení bodů
<code>poly</code>	: generuje polynom (<code>DOM_POLY</code>)
<code>polygon</code>	: generuje datovou strukturu <code>DOM_POLYGON</code> pro kreslení mnohoúhelníků
<code>rec</code>	: generuje rekurentní rovnici
<code>rectform</code>	: rozklad komplexního čísla na reálnou a imaginární část
<code>RootOf</code>	: symbolická množina kořenů polynomu (<code>RootOf</code>)
<code>series</code>	: rozvine výraz v řadu (generuje datový typ <code>Series::Puisseux</code> či <code>Series::gseries</code>)
<code>slot</code>	: definuje a/nebo čte atributy (slots) funkcí nebo doménových typů
<code>taylor</code>	: Taylorův rozvoj v bodě x_0 (<code>Series::Puisseux</code>)
<code>table</code>	: generuje prázdnou tabulku (<code>DOM_TABLE</code>)

Více konstrukcí speciálních datových struktur je definováno v knihovně `Dom`, například `Dom::Matrix()` generuje matice apod.

5. Operátory

Následující operátory mohou být užívány k volání systémových funkcí intuitivnějším, a pro mnohé srozumitelnějším, způsobem, např. `a+b` místo `_plus(a,b)`, `x<1` místo `_less(x,1)` atd.

operátor	priorita	systémová funkce	význam
::	2000	slot	přístup k atributům funkcí a doménových typů
,	1900	D	derivace (operátor diferencování)
[]	1800	_index	indexový operátor
.	1700	_concat	spojování řetězců
@@	1600	_fnest	opakování složení funkce
@	1500	_fconcat	složení dvou funkcí
!	1300	_fact	faktoriál
^	1200	_power	umocňování
*	1100	_mult	násobení
/	1100	_divide	dělení
+	1000	_plus	sčítání
-	1000	_subtract	odčítání
div	900	_div	celočíslné dělení
mod	900	_mod	funkce „modulo“
intersect	800	_intersect	průnik množin a intervalů
minus	700	_minus	rozdíl množin a intervalů
union	600	_union	sjednocení množin a intervalů
..	500	_range	obor hodnot
:=	500	_assign	přiřazení hodnoty
=	400	_equal	rovnost
<>	400	_unequal	nerovnost ($\neq$)
<	400	_less	porovnání
>	400	_less	porovnání
<=	400	_leequal	porovnání
>=	400	_leequal	porovnání
in	400	_in	vytvoří posloupnost prvků obsažených v objektu
\$	300	_seqgen	vytvoří posloupnost prvků
\$	300	_seqin	vytvoří posloupnost prvků
not	300	_not	logická „negace“
and	200	_and	logické „a“
or	100	_or	logické „nebo“
,		_exprseq	oddělovač mezi prvky seznamu (posloupnosti)
;		_stmtseq	oddělovač mezi příkazy; výpis výsledku není potlačen

:		<code>_stmtseq</code>	oddělovač mezi příkazy; výpis výsledku je potlačen
---	--	-----------------------	--

Uživatel si může definovat vlastní operátor pomocí příkazu `operator`.

6. Aritmetické funkce

Následující funkce mohou být používány pomocí operátorů nebo slovních příkazů, tj. `a+b` místo `_plus(a,b)`, `-a` místo `_negate(a)`, atd.

<code>_and</code>	:	logické „a“
<code>_assign</code>	:	přiřazovací funkce (přiřadí proměnné hodnotu)
<code>_break</code>	:	předčasné přerušení cyklu
<code>_case</code>	:	větvení v cyklu typu <code>switch</code>
<code>_concat</code>	:	spojování objektů
<code>_delete</code>	:	vymazání (odstranění) prvku nebo hodnoty
<code>_div</code>	:	celočíslné dělení
<code>_divide</code>	:	dělení
<code>_equal</code>	:	rovnost
<code>_fconcat</code>	:	složení funkcí
<code>_fnest</code>	:	opakované složení funkcí
<code>_for</code>	:	cyklus (vzestupný) typu „for“
<code>_for_down</code>	:	cyklus (sestupný) typu „for“
<code>_for_in</code>	:	cyklus přes prvky objektu
<code>_if</code>	:	podmíněný příkaz
<code>_index</code>	:	indexový operátor (přístup přes indexy)
<code>_intersect</code>	:	průnik množin a intervalů
<code>_leequal</code>	:	nerovnost $\leq$
<code>_lazy_and</code>	:	zkrácené vyhodnocování logického „a“
<code>_lazy_or</code>	:	zkrácené vyhodnocování logického „nebo“
<code>_less</code>	:	nerovnost $<$
<code>_minus</code>	:	rozdíl množin a intervalů
<code>_mod</code>	:	funkce „modulo“
<code>_mult</code>	:	násobení
<code>_negate</code>	:	negace nebo násobení -1
<code>_next</code>	:	přeskočí krok v cyklu
<code>_not</code>	:	logická negace
<code>_or</code>	:	logické „nebo“

<code>_plus</code>	: sčítání
<code>_power</code>	: umocnění
<code>_procdef</code>	: definice procedury
<code>_quit</code>	: ukončení <i>MuPADu</i> v UNIXu a LINUXu
<code>_range</code>	: rozsah
<code>_repeat</code>	: cyklus typu „repeat“
<code>_seqgen</code>	: vytvoří posloupnost prvků
<code>_seqin</code>	: vytvoří posloupnost prvků obsažených v objektu
<code>_subtract</code>	: odčítání
<code>_unequal</code>	: nerovnost typu $<>$ (různé, nerovno)
<code>_union</code>	: sjednocení množin a intervalů
<code>_while</code>	: cyklus typu „while“

7. Matematické symboly a konstanty

<code>C_</code>	: množina všech komplexních čísel $\mathbb{C}$
<code>CATALAN</code>	: konstanta: $\sum_{i=0}^{\infty} \frac{(-1)^i}{(2i+1)^2} = 0.915\,965\dots$
<code>complexInfinity</code>	: nevlastní bod komplexní roviny
<code>E</code>	: Eulerovo číslo: $\exp(1) = 2.718\,281\dots$
<code>EULER</code>	: Eulerova-Mascheroniova konstanta: $\gamma = \lim_{n \rightarrow \infty} \left(\sum_{i=1}^n \frac{1}{i} - \ln(n) \right) = 0.577\,215\dots$
<code>FAIL</code>	: chybový objekt, indikuje selhání výpočtu nebo neexistující prvek apod., jediný ob- jekt doménového typu <code>DOM_FAIL</code>
<code>FALSE</code>	: Booleova konstanta – nepravda
<code>I</code>	: imaginární jednotka
<code>infinity</code>	: nekonečno
<code>NIL</code>	: nulový objekt, jediný objekt doménového typu <code>DOM_NIL</code>
<code>PI</code>	: $\pi = 3.141\,592\,653\dots$
<code>Q_</code>	: množina všech racionálních čísel $\mathbb{Q}$
<code>R_</code>	: množina všech reálných čísel $\mathbb{R}$
<code>TRUE</code>	: Booleova konstanta – pravda
<code>undefined</code>	: nedefinovaná hodnota
<code>universe</code>	: teoretický soubor všech množin
<code>UNKNOWN</code>	: Booleova konstanta – neznámá hodnota
<code>Z_</code>	: množina všech celých čísel $\mathbb{Z}$

8. Matematické funkce

Následující matematické funkce jsou součástí *MuPADu* verze 2.0. Mohou být používány v argumentu funkcí typu `expand`, `float`, `simplify`, atd.

<code>abs</code>	: absolutní hodnota
<code>arccos</code>	: inverzní funkce k funkci <code>cos</code>
<code>arccosh</code>	: inverzní funkce k funkci <code>cosh</code>
<code>arccot</code>	: inverzní funkce k funkci <code>cot</code>
<code>arccoth</code>	: inverzní funkce k funkci <code>coth</code>
<code>arccsc</code>	: inverzní funkce k funkci <code>csc</code>
<code>arccsch</code>	: inverzní funkce k funkci <code>csch</code>
<code>arcsec</code>	: inverzní funkce k funkci <code>sec</code>
<code>arcsech</code>	: inverzní funkce k funkci <code>sech</code>
<code>arcsin</code>	: inverzní funkce k funkci <code>sin</code>
<code>arcsinh</code>	: inverzní funkce k funkci <code>sinh</code>
<code>arctan</code>	: inverzní funkce k funkci <code>tan</code>
<code>arctanh</code>	: inverzní funkce k funkci <code>tanh</code>
<code>arg</code>	: argument (polární úhel) komplexního čísla
<code>bernoulli</code>	: Bernoulliho čísla a polynomy
<code>besselI</code>	: modifikovaná Besselova funkce prvního druhu
<code>besselJ</code>	: Besselova funkce prvního druhu
<code>besselK</code>	: modifikovaná Besselova funkce druhého druhu
<code>besselY</code>	: Besselova funkce druhého druhu
<code>beta</code>	: Beta funkce, tj. $B(a,b) = \int_0^1 x^{a-1}(1-x)^{b-1} dx$
<code>binomial</code>	: binomický koeficient $\binom{m}{n}$
<code>ceil</code>	: nejmenší celé číslo $\geq x$
<code>Ci</code>	: $Ci(x) = \gamma + \ln(x) + \int_0^x (\cos(t) - 1)/t dt$
<code>cos</code>	: kosinus
<code>cosh</code>	: hyperbolický kosinus
<code>cot</code>	: cotangens
<code>coth</code>	: hyperbolický cotangens
<code>csc</code>	: $1/\sin(x)$
<code>csch</code>	: $1/\sinh(x)$
<code>dilog</code>	: dilogaritmická funkce $\int_0^x \ln(t)/(1-t) dt$
<code>dirac</code>	: Diracova δ -funkce
<code>Ei</code>	: $Ei(x) = \int_1^\infty t^{-1}e^{-tx} dt$
<code>erf</code>	: chybová funkce $\operatorname{erf}(x) = 2\pi^{-1/2} \int_0^x e^{-t^2} dt$
<code>erfc</code>	: doplněk chybové funkce, tj. $1 - \operatorname{erf}(x)$
<code>exp</code>	: exponenciální funkce e^x

fact	: faktoriál
floor	: největší celé číslo $\leq x$
frac	: desetinná část čísla
gamma	: Γ -funkce, tj. $\Gamma(a) = \int_0^\infty e^{-t} t^{a-1} dt$
heaviside	: primitivní funkce k Diracově δ -funkci
id	: identické zobrazení $x \rightarrow x$
igamma	: neúplná Γ -funkce $\text{igamma}(a, x) = \int_x^\infty e^{-t} t^{a-1} dt$, $\text{igamma}(a, 0) = \text{gamma}(a)$
Im	: imaginární část komplexního čísla
lambertV	: dolní větev Lambertovy funkce (řešení rovnice $ye^y = x$)
lambertW	: horní větev Lambertovy funkce
ln	: přirozený logaritmus
log	: logaritmus o libovolném základu, např. $\log_2 x$
polylog	: polylogaritmus, $Li_n(x) = \sum_{k=1}^\infty x^k / k^n$
psi	: polygamma funkce – derivace logaritmu Γ -funkce
Re	: reálná část komplexního čísla
round	: zaokrouhlení na nejbližší celé číslo
sec	: $\sec(x) = 1 / \cos(x)$
sech	: $\text{sech}(x) = 1 / \cosh(x)$
Si	: $\text{Si}(x) = \int_0^x t^{-1} \sin t dt$ („integral sinus“)
sign	: znaménko reálného nebo komplexního čísla
signIm	: znaménko imaginární části komplexního čísla
sin	: sinus
sinh	: hyperbolický sinus
sqrt	: druhá odmocnina
tan	: tangens (tg)
tanh	: hyperbolický tangens (tgh)
trunc	: odříznutí desetinné části čísla
zeta	: Riemannova ζ -funkce

9. Operace s polynomy

Následující funkce pracují s polynomy doménového typu `DOM_POLY` vytvořenými příkazem `poly`. Některé také pracují s polynomiálními výrazy (doménový typ `DOM_EXPR`).

coeff	: vrací nenulové koeficienty polynomu, např. <code>coeff(x^2-3*x+2)->{1,-3,2}</code>
--------------	--

<code>collect</code>	: vrátí koeficienty členů polynomu setříděného podle mocnin
<code>content</code>	: vrací největší společný dělitel všech koeficientů
<code>degree</code>	: vrací stupeň polynomu
<code>degreevec</code>	: vrací exponent(y) nejvyššího členu polynomu
<code>diff, D</code>	: spočte (parciální) derivace funkce, výrazu nebo polynomu
<code>divide</code>	: dělení polynomů se zbytkem
<code>evalp</code>	: hodnota polynomu v bodě
<code>factor</code>	: rozklad polynomu na součin dále nerozložitelných polynomů
<code>gcd</code>	: největší společný dělitel
<code>gcdex</code>	: největší společný dělitel polynomů (užitím rozšířeného euklidovského algoritmu)
<code>genpoly</code>	: generuje polynom pomocí b -adického rozvoje
<code>ground</code>	: absolutní člen polynomu
<code>irreducible</code>	: test nerozložitelnosti (irreducibility) polynomu
<code>lcm</code>	: nejmenší společný násobek polynomů
<code>lcoeff</code>	: koeficient u nejvyššího členu polynomu
<code>ldegree</code>	: stupeň členu s největším nenulovým exponentem
<code>lmonomial</code>	: nejvyšší netriviální monom polynomu
<code>lterm</code>	: nejvyšší člen polynomu
<code>mapcoeffs</code>	: aplikace funkce na koeficienty polynomu
<code>multcoeffs</code>	: násobení koeficientů polynomu skalárem
<code>norm</code>	: výpočet normy polynomu
<code>nterms</code>	: počet nenulových členů polynomu
<code>nthcoeff</code>	: n -tý nenulový koeficient polynomu
<code>nthmonomial</code>	: n -tý netriviální monom polynomu
<code>nthterm</code>	: n -tý nenulový člen polynomu
<code>pdivide</code>	: pseudo-dělitel polynomu v jedné proměnné
<code>poly</code>	: generátor (vytváření) polynomů
<code>polylib::decompose</code>	: funkcionální rozklad polynomu
<code>polylib::discrim</code>	: diskriminant polynomu

```

polylib::primpart : prvotní polynom (koeficienty vyděleny
                    největším společným násobkem)
polylib::resultant : determinant Sylvestrový matice
poly2list          : konvertuje polynom na seznam členů
tcoeff            : koeficient nejnižšího členu polynomu

```

Více procedur pracujících s polynomy je dostupných v knihovně `polylib`.

10. Operace s objekty *MuPADu*

```

alias              : definice zkratky (aliasu) pro objekt
anames             : seznam identifikátorů, které mají
                    přiřazenou hodnotu nebo vlastnost
append            : připojení prvku k seznamu
assign            : přiřazení hodnoty dané jako rovnice
assignElements    : přiřazení hodnoty pro prvky pole,
                    seznamu nebo tabulky
assume            : přidělí vlastnosti identifikátoru
bool              : logická hodnota výrazu s výstupem typu
                    TRUE nebo FALSE
combine           : úprava výrazu podle zadaných pravidel
                    (opak expand)
contains          : test na existenci prvku v seznamech,
                    množinách a tabulkách
delete            : odstraní hodnotu identifikátoru
denom             : jmenovatel racionálního výrazu
domain           : definuje nový datový (doménový) typ
                    (synonym pro datový typ – doménu)
domtype          : zjištění doménového typu objektu
expand            : rozklad výrazu na „základní prvky“ (opak
                    combine); například
                    expand((x+y)^2)=2xy+x^2+y^2
export            : zpřístupnění dané knihovny funkcí
expose            : zobrazí zdrojový MuPADovský kód funkcí
                    z knihoven a definice doménových typů
expr             : převod (konverze) objektu na prvky
                    základní datové domény
extnops           : vrací počet operandů interní reprezentace
                    objektu v doménovém typu

```

<code>genident</code>	: vytvoření nového, dosud nepoužitého, identifikátoru
<code>getprop</code>	: vrací matematické vlastnosti výrazu, např. zda se jedná o celé číslo apod.
<code>has</code>	: testuje, zda se daný objekt vyskytuje v syntaxi některého jiného objektu
<code>hastype</code>	: testuje, zda se daný typ objektu vyskytuje v syntaxi některého jiného objektu
<code>history</code>	: přístup do tabulky historie <i>MuPADu</i>
<code>indets</code>	: vrací neurčené proměnné a konstanty vyskytující se ve výrazu
<code>is</code>	: kontrola matematických vlastností výrazu, např. zda se jedná o celé číslo apod.
<code>lasterror</code>	: znovu zobrazí poslední chybovou hlášku
<code>length</code>	: vrací délku (velikost) objektu <i>MuPADu</i>
<code>lhs</code>	: levá strana výrazu
<code>listlib::insert</code>	: vložení prvku do seznamu
<code>map</code>	: aplikuje funkci na všechny operandy objektu
<code>maprat</code>	: aplikuje funkci na „ <i>racionalizovaný</i> “ výraz
<code>match</code>	: porovnává syntakticky vzor s výrazy
<code>nops</code>	: vrací počet operandů objektu
<code>numer</code>	: čítec racionálního výrazu
<code>op</code>	: vrací operandy objektu
<code>protect</code>	: ochrana identifikátoru proti zápisu
<code>radsimp</code>	: zjednodušení výrazů s odmocninami
<code>rationalize</code>	: převod výrazu na „ <i>racionální</i> “ tvar
<code>rewrite</code>	: vyjádří výraz v matematicky ekvivalentní formě (dle „potřeb uživatele“)
<code>rhs</code>	: pravá strana výrazu
<code>select</code>	: vybírá prvky objektů splňující dané podmínky
<code>simplify</code>	: zjednodušení výrazu
<code>slot</code>	: vrací a/nebo čte atributy (slots) funkcí nebo doménových typů
<code>split</code>	: rozdělení prvků (operandů) objektu podle zadané vlastnosti
<code>subs</code>	: substituce do objektu
<code>subsex</code>	: rozšířená substituce (složitých výrazů)

<code>subsop</code>	: substituce operandů
<code>sysorder</code>	: porovnání objektů na základě jejich vnitřního uspořádání v <i>MuPADu</i>
<code>testtype</code>	: testování objektu na zadaný datový typ
<code>traperror</code>	: zachycení chyb
<code>type</code>	: zobrazí typ objektu
<code>unalias</code>	: odstraní zkratku (alias)
<code>unassume</code>	: odstraní vlastnosti identifikátoru
<code>unexport</code>	: znepřístupnění dané knihovny funkcí
<code>unprotect</code>	: odstraňuje ochranu identifikátorů proti zápisu
<code>zip</code>	: sloučení seznamů pomocí dané funkce

11. Manipulace s textovými řetězci

<code>concat</code>	: spojování řetězců
<code>.</code>	: spojování řetězců
<code>expr2text</code>	: převod výrazu na řetězec
<code>ftextinput</code>	: čtení textového souboru
<code>int2text</code>	: konverze celého čísla na řetězec
<code>length</code>	: zjištění délky řetězce
<code>revert</code>	: „obrácení“ řetězce, např. <code>revert("abc")="cba"</code>
<code>strmatch</code>	: porovnání řetězců
<code>substring</code>	: výběr podřetězců
<code>tbl2text</code>	: konverze tabulky na řetězec
<code>text2expr</code>	: konverze řetězce na výraz
<code>text2int</code>	: konverze řetězce na celé číslo
<code>text2list</code>	: konverze řetězce na seznam
<code>text2tbl</code>	: konverze řetězce na tabulku
<code>textinput</code>	: interaktivní vstup textu (přes prompt)

Více procedur pro manipulaci s textovými řetězci naleznete v knihovně `stringlib`.

12. Procedury pro vstup a výstup

<code>fclose</code>	: uzavření souboru
---------------------	--------------------

```

finput      : čtení MuPADovského objektu ze souboru
fopen       : otevření souboru pro čtení a pro zápis
fprint      : zápis dat do souboru
fread       : čtení a provedení MuPADovského souboru
ftextinput  : vstup ze souboru textového tvaru
input       : interaktivní vstup hodnot (přes prompt)
print       : výstup (jakéhokoliv) objektu na obrazovku
protocol    : otevření a uzavření protokolu o práci
read        : čtení ze souboru
textinput   : interaktivní vstup textu (přes prompt)
write       : uložení hodnot proměnných do souboru

```

Více procedur pro řízení vstupů a výstupů naleznete v knihovnách `import` a `output`.

13. Procedury pro grafiku

```

plot        : vykreslí jednorozměrný grafický objekt
plot2d      : vykreslí dvourozměrný grafický objekt
plot3d      : vykreslí třírozměrný grafický objekt
plotfunc2d  : vykreslí graf funkce jedné proměnné
plotfunc3d  : vykreslí graf funkce dvou proměnných

```

Mnohem více procedur pro grafiku naleznete v knihovně `plot`.

14. Procedury pro řízení výpočtů

```

bool        : vyhodnocení logických výrazů s výstupem
              TRUE nebo FALSE
context     : vyhodnocení procedury (v kontextu
              volající procedury)
eval        : vyhodnocení objektu
evalassign  : vyhodnocení výrazu do dané hloubky
              a jeho přiřazení
freeze      : vytvoření neaktivní kopie funkce
hold        : zabránění vyhodnocení objektu
indexval    : indexovaný přístup k vektorům a maticím
              bez jejich vyhodnocení

```


<code>last</code>	: přístup k předchozímu výsledku
<code>%</code>	: přístup k předchozímu výsledku (lze užít místo <code>last</code>)
<code>level</code>	: vyhodnocení objektu do předem dané hloubky
<code>unfreeze</code>	: znovu aktivuje činnost funkce
<code>val</code>	: nahradí všechny identifikátory objektu jejich hodnotami (ekvivalentní k <code>level(.,1)</code> , ale bez vnitřního zjednodušení)

15. Příkazy pro definování procedur

<code>_procdef</code>	: definice procedury
<code>args</code>	: přístup k argumentům procedury
<code>error</code>	: výstup chybového hlášení
<code>return</code>	: ukončení procedury a vrácení výsledků
<code>testargs</code>	: kontrola datových typů argumentů na vstupu

16. Matematické funkce a algoritmy

S některými funkcemi uvedenými v této části se pozorný čtenář mohl setkat v sekcích jiných. Zde je znovu uvádíme pro úplnost a usnadnění hledání v tomto textu.

<code>abs</code>	: absolutní hodnota
<code>arg</code>	: argument (polární úhel) komplexního čísla
<code>asympt</code>	: počítá asymptotický rozvoj výrazu
<code>binomial</code>	: kombinační číslo $\binom{m}{n}$
<code>ceil</code>	: nejmenší celé číslo $\geq x$
<code>conjugate</code>	: komplexně sdružené číslo
<code>D</code>	: derivace (operátor diferencování)
<code>'</code>	: derivace (operátor diferencování)
<code>diff</code>	: derivování výrazu nebo polynomu
<code>discont</code>	: bod(y) nespojitosti funkce
<code>div</code>	: celá část racionálního výrazu
<code>factor</code>	: faktorizace polynomů, výrazů a celých čísel
<code>float</code>	: převod výrazu na desetinné číslo

<code>floor</code>	: největší celé číslo $\leq x$
<code>frac</code>	: desetinná část čísla
<code>gcd</code>	: největší společný dělitel
<code>gcdex</code>	: největší společný dělitel (použitím rozšířeného euklidovského algoritmu)
<code>ifactor</code>	: rozklad přirozeného čísla na prvočinitele
<code>igcd</code>	: největší společný dělitel celých čísel
<code>igcdex</code>	: největší společný dělitel celých čísel (rozšířený Euklidův algoritmus)
<code>ilcm</code>	: nejmenší společný násobek celých čísel
<code>Im</code>	: imaginární část komplexního čísla
<code>int</code>	: integrace (výpočet určitých a neurčitých integrálů)
<code>intersect</code>	: průnik množin a intervalů
<code>isprime</code>	: test prvočíselnosti
<code>isqrt</code>	: celočíselná aproximace odmocniny z celého čísla; vrací <code>floor(sqrt(x))</code>
<code>iszero</code>	: test nulovosti
<code>ithprime</code>	: vrací i -té prvočíslo
<code>lcm</code>	: nejmenší společný násobek polynomů
<code>limit</code>	: výpočet hodnoty limity
<code>linsolve</code>	: řešení soustav lineárních rovnic
<code>lllint</code>	: výpočet LLL-redukované báze svazu
<code>max</code>	: maximum
<code>min</code>	: minimum
<code>minus</code>	: rozdíl množin a intervalů
<code>mod, modp, mods</code>	: funkce „modulo“
<code>nextprime</code>	: nejmenší prvočíslo $\geq x$, $x \in \mathbb{N}$
<code>norm</code>	: výpočet normy polynomů, vektorů a matic
<code>normal</code>	: převod racionálních výrazů na základní tvar (normalizace objektů)
<code>not</code>	: logická negace
<code>numeric::fft</code>	: rychlá Fourierova transformace
<code>numeric::invfft</code>	: inverzní rychlá Fourierova transformace
<code>numeric::lagrange</code>	: interpolace pomocí polynomů
<code>numeric::cubicSpline</code>	: interpolace kubickým splinem
<code>numlib::contfrac</code>	: rozvoj ve tvaru řetězového zlomku
<code>or</code>	: logické „nebo“
<code>pade</code>	: Padeho aproximace

<code>partfrac</code>	: rozklad na parciální zlomky
<code>polylib::sqrfree</code>	: rozklad polynomu na součin mocnin
<code>powermod</code>	: určí efektivně $b^p \bmod m$
<code>product</code>	: součin
<code>random</code>	: generátor náhodných čísel
<code>Re</code>	: reálná část komplexního čísla
<code>rectform</code>	: rozklad komplexního čísla na reálnou a imaginární část
<code>revert</code>	: reverze (obrácení) řad, řetězců a seznamů
<code>round</code>	: zaokrouhlení na nejbližší celé číslo
<code>series</code>	: rozvoj ve tvaru řady
<code>sign</code>	: znaménko reálného čísla (pro komplexní číslo $\text{sign}(z) = z / \text{abs}(z)$)
<code>signIm</code>	: znaménko imaginární části komplexního čísla
<code>solve</code>	: řeší systémy rovnic a nerovnic
<code>solveLib::pdioe</code>	: řešení polynomiálních diofantických rovníc
<code>sort</code>	: třídění seznamů
<code>sum</code>	: výpočet formálních i numerických součtů
<code>taylor</code>	: Taylorův rozvoj okolo daného bodu
<code>trunc</code>	: odříznutí všech číslic za desetinnou tečkou
<code>union</code>	: sjednocení množin

Více speciálních matematických funkcí naleznete v knihovnách popsaných v oddíle 17..

17. Knihovny a moduly

Jednotlivé knihovny *MuPADu* 2.0 obsahují daleko více speciálních matematických funkcí než zde uvádíme. Lze přitom očekávat, že jednotlivé knihovny budou neustále rozšiřovány. Přehled současných funkcí se volá příkazem `info()` nebo `?`. Nápovědu k jednotlivé funkci dostanete příkazem

`?libraryname::functionname`. Nápovědu k modulu získáte pomocí příkazu `modulename::doc()`. Následující tabulka shrnuje nejdůležitější knihovny a moduly *MuPADu*.

<code>adt</code>	: abstraktní datové typy
<code>Ax</code>	: předdefinované axiomy; generátor axiomů
<code>Cat</code>	: předdefinované kategorie; generátor kategorií
<code>combinat</code>	: kombinatorika
<code>dertools</code>	: diferenciální rovnice
<code>Dom</code>	: předdefinované datové struktury a domény
<code>fp</code>	: funkcionální programování
<code>generate</code>	: generování výstupů ve formátech C, Fortran a $\text{\TeX}$
<code>groebner</code>	: výpočet Groebnerových bází (ideálů pro polynomy)
<code>import</code>	: import dat
<code>intlib</code>	: symbolické a numerické integrování
<code>linalg</code>	: lineární algebra
<code>linopt</code>	: lineární optimalizace (lineární programování)
<code>listlib</code>	: manipulace se seznamy
<code>matchlib</code>	: porovnávání vzoru s výrazy
<code>misc</code>	: pomocné funkce (miscelanea)
<code>module</code>	: správa dynamických modulů
<code>Network</code>	: teorie grafů
<code>numeric</code>	: algoritmy numerické matematiky
<code>numlib</code>	: teorie čísel
<code>orthpoly</code>	: ortogonální polynomy
<code>output</code>	: formátování výstupu
<code>plot</code>	: procedury pro kreslení
<code>polylib</code>	: procedury pro manipulaci s polynomy
<code>Pref</code>	: uživatelská nastavení systémového prostředí
<code>prog</code>	: programování a trasování
<code>property</code>	: vlastnosti identifikátorů a výrazů

RGB	: předdefinované názvy barev (je použit aditivní způsob skládání barev <i>červená–zelená–modrá</i>)
solve	: řešení, nejenom lineárních, rovnic
stats	: statistické funkce
stdlib	: standardní knihovna funkcí
stringlib	: manipulace s textovými řetězci
student	: některé elementární algoritmy vhodné pro výuku matematiky
transform	: integrální transformace
Type	: objekty pro syntaktickou kontrolu datových typů

18. Trasování, ladění a měření času

Procedura **debug** pro trasování není k dispozici u verze *Light*, nicméně je k dispozici ve volně šířené verzi pro operační systém *Linux*.

debug	: procedura pro trasování (ladění) procedur a funkcí
prog::profile	: statistika o době potřebné k provedení procedur a funkcí
prog::trace	: protokol o provedení procedury nebo funkce
rtime	: celková doba výpočtu (v milisekundách)
time	: čas použití procesoru

Více procedur pro ladění naleznete v knihovně **prog**.

19. Technické parametry

bytes	: obsazení paměti
export	: zpřístupnění knihovny nebo funkce z knihovny
external	: externí přístup k modulům
getpid	: v UNIXu a LINUXu vrací ID číslo procesu jádra <i>MuPADu</i>
loadlib	: načtení knihovny (doporučujeme raději používat příkaz package)
loadmod	: načtení dynamického modulu
loadproc	: načtení procedury z knihovny
operator	: definuje nový operátor

<code>package</code>	: načtení uživatelem definované knihovny
<code>patchlevel</code>	: informace o nainstalovaných opravách knihoven <i>MuPADu</i>
<code>pathname</code>	: upravuje zápis cest k souborům podle typu operačního systému
<code>quit</code>	: ukončení <i>MuPADu</i> pod UNIXem a LINUXem
<code>register</code>	: registrace <i>MuPADu</i>
<code>reset</code>	: resetování (inicializace) prostředí <i>MuPADu</i>
<code>sysname</code>	: vrací název užívaného operačního systému
<code>system</code>	: spuštění příkazu operačního systému
<code>unexport</code>	: znepřístupnění dané knihovny funkcí
<code>unloadmod</code>	: odstranění načtených modulů
<code>version</code>	: informace o číslu verze <i>MuPADu</i>

20. Novinky a změny v *MuPADu* verze 2.5

MuPAD 2.5 obsahuje 70 nových funkcí v knihovně `stats`, kompletně přepracovanou knihovnu `combinat` obsahující nyní 11 podknihoven a 4 nové funkce v knihovně `numeric`. Novinkou je možnost přímého propojení *MuPAD* u s numerickým systémem *Scilab*. Další informace naleznete v dokumentu *News and Changes*, který je obsažen v distribuci *MuPADu* 2.5. Nové funkce, klíčová slova, operátory a datové typy:

<code>==></code>	: logická implikace
<code><==></code>	: logická ekvivalence
<code>...</code>	: zápis intervalových výrazů typu <code>DOM_INTERVAL</code>
<code>assert</code>	: definuje tvrzení pro debugging
<code>card</code>	: počet prvků množiny
<code>copyClosure</code>	: kopie procedury s ochranou lokálních identifikátorů
<code>DOM_INTERVAL</code>	: obdélníkové komplexní intervaly
<code>FILEPATH</code>	: cesta k načtenému souboru
<code>fname</code>	: název otevřeného souboru
<code>frame</code>	: změni rámec platnosti lokálních identifikátorů
<code>frandom</code>	: generuje náhodné číslo z intervalu $[0, 1)$
<code>hull</code>	: konvertuje numerické a intervalové výrazy na typ <code>DOM_INTERVAL</code>
<code>hypergeom</code>	: hypergeometrická funkce
<code>interval</code>	: konvertuje konstantní podvýrazy na intervaly
<code>PACKAGEPATH</code>	: cesta k balíkům knihoven funkcí
<code>save</code>	: chrání globální identifikátory v proceduře

(eds.) Jaromír Antoch, Daniel Hlubinka a Ivan Saxl

PRAVDĚPODOBNOST A STATISTIKA NA STŘEDNÍ ŠKOLE

Vydal

MATFYZPRESS

vydavatelství

Matematicko-fyzikální fakulty

Univerzity Karlovy v Praze

Sokolovská 83

186 75 Praha 8

jako svou 125 publikaci

Z připravených předloh v systému $\text{T}_{\text{E}}\text{X}$

vytisklo Repro středisko UK MFF

Sokolovská 83, 186 75 Praha 8

Vydání první

Praha 2005

ISBN 80-86732-23-1