

Vysoká škola ekonomická v Praze

R SNADNO A RYCHLE 1

ÚVOD DO JAZYKA



Jakub Danko, Karel Šafr

2020



OECONOMICA
Nakladatelství VŠE

Autoři:

Ing. Jakub Danko, Ph.D. (Vysoká škola ekonomická v Praze)

Ing. Karel Šafr, Ph.D. (Vysoká škola ekonomická v Praze)

Recenzenti:

Mgr. Vítězslav Štembera, Ph.D. (Technická univerzita Vídeň)

Ing. Zdeněk Šulc, Ph.D. (Vysoká škola ekonomická v Praze)

Podakovanie

Na tomto mieste by sme sa radi poďakovali najmä doc. RNDr. Ivane Malej, CSc. (Vysoká škola ekonomická v Prahe) za veľké množstvo času, ktoré venovala pomoci pri písaní tejto našej „prvotiny“. Bez jej odborných pripomienok, komentárov a rád by táto kniha určite nevznikla v súčasnom formáte, rozsahu a kvalite. Okrem toho by sme chceli poďakovať aj recenzentom tejto knihy, ktorými sú: Mgr. Vítězslav Štembera, Ph.D. (Technická univerzita Viedeň) a Ing. Zdeněk Šulc, Ph.D. (Vysoká škola ekonomická v Prahe), za množstvo kvalitných pripomienok, ktoré pomohli k dokončeniu tejto knihy. Rovnako si vážime aj to, že venovali množstvo svojho voľného času detailnému recenzovaniu tejto knihy. Ďalej by sme radi poďakovali Ing. Veronike Ptáčkovej a Mgr. Beáte Ráczovej, ktoré vykonali prvotnú gramatickú a štylistickú korektúru tejto knihy. Vďaka, samozrejme, patrí aj našim priateľom a rodine, ktorí nás pri písaní tejto knihy neustále podporovali a motivovali k jej dokončeniu.

V neposlednom rade je potrebné povedať, že vznik knihy bol podporený projektom internej rozvojovej súťaže Vysokej školy ekonomickej v Prahe (IRS) F4/18/2019, za ktorého finančnú a materiálnu podporu sa chceme taktiež poďakovať.

Autori

Obsah

1. Predstavenie programovacieho jazyka R	6
1.1 Výhody a nevýhody práce v R	7
1.2 Inštalácia R	9
1.2.1 Grafické užívateľské rozhranie verzus editor	9
1.3 Inštalácia RStudio	10
1.4 RStudio a základy práce v tomto prostredí	10
1.5 R ako kalkulačka – práca v konzole	12
1.6 Pomoc a dokumentácia	14
1.6.1 Dokumentácia k funkciám	14
1.6.2 Dokumentácia k balíkom	17
1.6.3 Demonštračné kódy	17
1.6.4 Dokumentácia ku konkrétnej téme	18
1.6.5 Vinety	18
1.6.6 Dokumentácia z externých zdrojov	19
1.7 Najčastejšie chyby pri práci s R	20
1.8 Práca s balíkmi	21
1.9 Vytvorenie pracovného adresára	24
1.10 Vytvorenie skriptu	25
2. Vytváranie objektov	27
2.1 Pravidlá pri vytváraní objektov	30
3. Dátové typy	31
3.1 Zmena dátového typu	33
4. Operátory	36
4.1 Aritmetické operátory	36
4.2 Relačné (porovnávacie) operátory	37
4.3 Logické operátory	38
4.4 Operátory, prostredníctvom ktorých vyberáme časti dátových štruktúr	40
4.5 Ďalšie používané operátory	40

5. Funkcie	41
6. Dátové štruktúry	46
6.1 Homogénne dátové štruktúry	47
6.1.1 Atomický vektor	47
6.1.2 Atomická matica	68
6.1.3 Atomické pole	104
6.2 Nehomogénne dátové štruktúry	108
6.2.1 Zoznam (list)	108
6.2.2 Data frame	115
6.3 Zistenie obsahu dátovej štruktúry	119
7. Základy spracovania dát	122
7.1 Import dát	122
7.1.1 Funkcie vhodné na import dát a ich základné parametre	124
7.1.2 Načítavanie dát z internetu	126
7.1.3 Načítavanie dát z pracovného adresára	127
7.2 Export dát z R do štandardných formátov	133
7.3 Export dát z R do formátu RData	135
7.3.1 Funkcie saveRDS() a readRDS()	135
7.3.2 Funkcie save() a load()	136
7.3.3 Funkcie save.image() a load()	137
7.4 Práca s chýbajúcimi pozorovaniami v R	137
8. Práca s dátovými súbormi	141
8.1 Balík datasets a jeho vybrané datasety	142
8.2 Praktické spracovanie dát z datasetu mtcars	143
Zoznam použitej literatúry	159

Kapitola 1

Predstavenie programovacieho jazyka R

R predstavuje voľne šíriteľný programovací jazyk, ktorý je veľmi vhodný a využívaný na štatistickú analýzu, spracovanie a vizualizáciu údajov. Z programátorského pohľadu ide o skriptovací jazyk¹. Výhodou skriptovacích jazykov je ľahšia údržba a vývoj. Skript je možné kedykoľvek spustiť bez nutnosti ho vždy znovu skompilovať. Programátorovi v tomto prípade stačí obyčajný textový editor, nepotrebuje teda žiadny špeciálny kompilátor alebo vývojový nástroj. Okrem štatistiky sa **R** využíva aj v množstve iných oblastí, napríklad v sociálnych vedách alebo v biológii. Napríklad vo financiách a bankovníctve sa používa pri detekcii a identifikácii podvodov, v bioinformatike pri vývoji liečiv alebo v analýze sociálnych médií pri vyhľadávaní potenciálnych zákazníkov pre online cielenie reklamných kampaní. Univerzálnosť a rozsah použitia tohto jazyka rastie veľmi rýchlo, vzhľadom na to, že neustále vznikajú nové knižnice príkazov vhodné na aplikáciu v konkrétnej oblasti. Je to dané jeho obrovskou výhodou, keďže ide o takzvaný *open source* programovací jazyk, ktorý je pre používateľov zdarma, čo znamená, že nielen jeho používateľom, ale aj vývojárom sa môže stať prakticky ktokoľvek. V dôsledku vyššie uvedených dôvodov využívajú **R** v svojej analytickej práci mnohé renomované spoločnosti ako napríklad Facebook, Google, LinkedIn, IBM či Twitter. Vzhľadom na jeho neustále rastúci trend použitia v praxi, akademickej sfére a postupného zapájania do vyučovacieho procesu u väčšiny štatistických predmetov vyučovaných na Vysokej škole ekonomickej v Prahe sme sa rozhodli pre zostavenie študijného materiálu, v ktorom čitateľov zoznámime s týmto programovacím jazykom a uvidíme manuál, ako s ním krok za krokom pracovať.

¹Skriptovací jazyk je počítačový programovací jazyk pôvodne navrhnutý na uľahčenie (zautomatizovanie) operácií v počítači. Program v skriptovacom jazyku sa nazýva skript. Pre viac informácií pozri Kapitulu 1.10.

1.1 Výhody a nevýhody práce v R

Ako bolo uvedené v úvode, **R** je v súčasnosti veľmi obľúbený programovací jazyk a v tejto podkapitole si uvedieme jeho základné výhody a, samozrejme, aj jeho nevýhody pri porovnaní s konkurenčnými produktmi.

- Výhody práce v **R**:

- **Otvorený zdrojový kód** znamená, že s týmto programom môže pracovať každý bez toho, aby k tomu potreboval nejakú licenciu alebo zaplatenie poplatku. Ďalej môže ktokoľvek prispieť k rozvoju **R** prispôbením jeho balíkov (knižníc, pre viac informácií pozri Kapitulu 1.8), vývojom nových a riešením špecifických problémov.
- Práca v tomto programe je veľmi vhodná pre **spracovanie** veľkého množstva chaoticky usporiadaných **dát** a ich **transformáciu** do štruktúrovanej formy, na čo sú veľmi vhodné balíky ako napríklad **dplyr** alebo **readr**.
- Program má dostupnú **obrovskú kolekciu balíkov** (packages), ktorým je podrobnejšie venovaná Kapitola 1.8, ktorých počet neustále rastie (na začiatku roka 2020 je na centrálnom úložisku balíkov CRAN ([odkaz tu](#)) vyše 15-tisíc dostupných balíkov), čo svedčí o rastúcom trende a rozsahu použitia programovacieho jazyka v rôznych oblastiach.
- Prostredie programovacieho jazyka ponúka vytváranie veľmi kvalitných a estetických grafov. **Vizualizácia** s využitím napríklad balíkov **ggplot2** alebo **plotly** vytvára profesionálne vyzerajúce grafy s možnosťou rôznych úprav a prispôbení, ktoré sa odlišujú od iných programovacích jazykov.
- Programovací jazyk je **vysoko kompatibilný** s mnohými inými programovacími jazykmi ako C, C++, Java alebo Python a môže byť integrovaný napríklad s technológiou Hadoop, ktorá umožňuje uloženie údajov na veľké množstvo samostatných počítačov.
- Použitie programovacieho jazyka nezávisí od platformy, je **kompatibilný so všetkými najčastejšie používanými operačnými systémami**, čo znamená, že je možné ho používať ako v prostredí Windows, tak v prostredí Linux aj v prostredí Mac.
- Okrem kompatibility s inými programovacími jazykmi a operačnými systémami, je tento programovací jazyk flexibilný a dokáže pracovať aj s obrovským množstvom **súborových formátov**. Okrem klasických dát vo formáte .csv, .txt, .xls či .xlsx dokáže pracovať aj s datasetmi komerčných programov ako SAS, SPSS, STATA a podobne.
- Prostredníctvom balíkov ako **Shiny** a **Markdown** je možné veľmi jednoducho vytvárať komplexné **správy**, **dokumenty** a kompletizovať výsledky analýz získaných z programu priamo v užívateľskom prostredí. V týchto správach je

možné priamo pracovať s údajmi, grafmi a skriptmi, ktoré sú v nich zabudované. Okrem toho je možné vytvorenie **interaktívnych webových aplikácií**, v ktorých sa výsledky menia okamžite so zmenou vstupných údajov, čo má veľké možnosti použitia ako v praxi, tak aj vo vyučovacom procese.

- Programovací jazyk **R** je v prostredí štatistickej komunity považovaný za najviac rozšírený (tzv. **lingua franca štatistiky**), to je hlavný dôvod, prečo je dominantným medzi ostatnými programovacími jazykmi pri vývoji štatistických nástrojov.
 - Vďaka rozšíreniu tohto programovacieho jazyka vzniká obrovská **komunita** vývojárov, používateľov a programátorov, ktorí sú ochotní pomôcť a zdieľať svoje znalosti s ostatnými. Je preto možné vyhľadať pomoc na rôznych blogoch a iných on-line platformách na to určených. Pre programátorov asi najznámejšiu komunitu predstavuje Stack Overflow ([odkaz tu](#)), kde používatelia môžu nájsť tipy a odpovede týkajúce sa programovania v rôznych programovacích jazykoch. Stránka **R** blogerov ([odkaz tu](#)) je veľmi vhodná pre začínajúcich používateľov, nakoľko sa na nej nachádza množstvo návodov na vykonanie rôznych analýz spolu s príkladmi a vzorovými riešeniami na vybraných dátových súboroch. Používateľom, ktorí majú radi interaktívnu výuku formou videí, odporúčame kurzy na stránkach DataCamp ([odkaz tu](#)) či Coursera ([odkaz tu](#)).
- Nevýhody práce v **R**:
 - Jednou z nevýhod tohto programovacieho jazyka je fakt, že vychádza z programovacieho jazyka **S**, čo znamená, že v základnom tvare nepodporuje dynamickú ani 3D grafiku. Túto nevýhodu však možno kompenzovať použitím vyššie uvedených balíkov (napríklad **ggplot2**) vhodných na pokročilú vizualizáciu.
 - Z pohľadu spracovania údajov môže nastávať problém pri využívaní pamäte. V prostredí **R** sa objekty ukladajú do fyzickej pamäte, čo je rozdiel napríklad v porovnaní s programom Python, ktorý využíva menej pamäte. **R** taktiež vyžaduje všetky údaje na jednom mieste, preto nie je veľmi vhodný pri spracovaní tzv. veľkých dát (big data). Tento nedostatok je možné riešiť opäť prostredníctvom balíkov, ktoré sú určené na spracovanie veľkých databáz, prípadne integráciou s technológiou Hadoop uvedenou vyššie.
 - V porovnaní so spomínaným programom Python v programovacom jazyku **R** absentujú niektoré prvky zabezpečenia, preto sa až tak často nevyužíva vo webových aplikáciách.
 - Programovací jazyk **R** nie je úplne jednoduchý na učenie a práca v tomto prostredí nie je úplne intuitívna a jednoduchá (user friendly). Preto sa niekedy môže zdať ľuďom, ktorí nemajú predchádzajúce skúsenosti s programovaním, práca v **R** náročná. Je pravda, že existujú jednoduchšie klikacie programy, ktoré môžu používateľom pripadať viac intuitívne, tie sú však často komerčné,

prípadne majú obmedzenú funkcionálnosť. Práve z tohto dôvodu sa najmä u začínajúcich používateľov odporúča využívať nejaký editor, ktorý prácu zjednodušuje. Jedným z veľmi vhodných editorov je **RStudio**, ktoré je rovnako ako samotný programovací jazyk **R** k dispozícii zadarmo na nekomerčné použitie.

- V porovnaní s programami **MATLAB** a **Python** je práca s balíkmi a samotné programovanie v **R** pomalšie. Tieto nedostatky sa dajú riešiť využitím na to vhodných funkcií, napríklad v prípade práce s maticami využívať radšej funkcie typu `apply()` ako cykly a podobne. Treba však povedať, že pre začiatočníka aj základného používateľa tohto jazyka je jeho rýchlosť dostatočná a v prípade aplikácie na štandardné typy analýz úplne postačujúca.
- Vzhľadom na to, že podieľať sa na vývoji môže ktokoľvek, aj napriek viacnásobnej kontrole z pohľadu komunity používateľov môže byť kvalita niektorých (najmä najnovších a málo preverených) balíkov nepostačujúca, pričom tu tak tiež abscentuje nejaké inštitucionalizované pracovisko zákazníckej podpory, kde by bolo možné sťažovať sa, ak niečo nefunguje.

1.2 Inštalácia R

Ešte predtým, ako začneme pracovať v **R** a vytvoríme si napríklad svoj prvý príkaz či nadefinujeme prvú funkciu, je potrebné si program **R** nainštalovať. Inštalácia je dostupná na oficiálnej stránke tohto programu ([odkaz tu](#)). Pri inštalácii je potrebné zvoliť si operačný systém (Linux, Mac alebo Windows), stiahnuť si inštalačný súbor a postupovať pri inštalácii podľa pokynov na obrazovke. Po úspešnej inštalácii je možné spustiť nainštalovaný program prostredníctvom grafického užívateľského rozhrania (GUI = graphical user interface). Vzhľad tohto rozhrania sa líši v závislosti od verzie **R** a operačného systému. V prípade operačného systému Linux toto rozhranie viac pripomína príkazový riadok a je aj používateľsky jednoduchšie.

1.2.1 Grafické užívateľské rozhranie verzus editor

Práca v grafických užívateľských rozhraniach nie je najmä pre začínajúcich používateľov jednoduchá a vyžaduje si aspoň základnú znalosť programovania a syntaxe jazyka. Preto odporúčame pri práci v **R** používať editor **RStudio**, ktorý predstavuje integrované vývojové prostredie (IDE = integrated development environment) pre **R**. V tomto prostredí sa práca výrazne uľahčuje, nakoľko v porovnaní s GUI umožňuje pracovať oveľa efektívnejšie. Napríklad pri práci s viacerými zátvorkami sú tieto zátvorky farebne odlíšené, ako to poznáme z prostredia MS Excel pre lepšiu prehľadnosť. Príkazy, ktoré v prostredí zapisujeme, sa automaticky dopĺňajú, aby sme nemuseli všetko vpisovať a ušetrili si tým prácu. Prostredie ponúka jednoduchšiu správu balíkov a prácu s grafickými výstupmi. Funkcionalita je rozšírená napríklad o tvorbu dokumentov (**R Markdown**) alebo interak-

tívných webových aplikácií (Shiny). Okrem toho je aj orientácia v tomto prostredí oveľa jednoduchšia a intuitívnejšia.

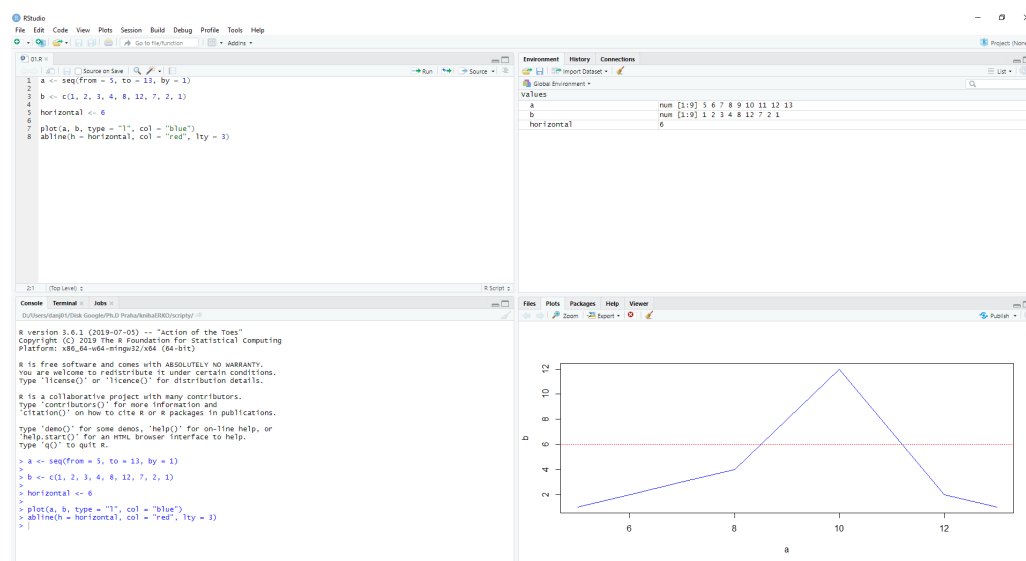
V prípade, ak vám toto prostredie nevyhovuje, existujú aj iné textové editory, s ktorými môžete pracovať. Príkladom sú editory **Rattle** (vhodný najmä na data mining), **StatET**, **ESS**, **Tinn-R** a podobne.

1.3 Inštalácia RStudio

Inštalácia **RStudio** je dostupná na oficiálnej stránke tohto editora ([odkaz tu](#)). Pre potreby bežného používateľa postačí otvorená licencia verzie **RStudio Desktop** (Open Source Licence), ktorá je k dispozícii zadarmo. Komerčné využitie tohto editora je spoplatnené. Po kliknutí na stiahnutie editora (download) máme opäť v ponuke verzie pre všetky najrozšírenejšie druhy operačných systémov. Po stiahnutí spustíme inštaláčny súbor a postupujeme podľa pokynov na obrazovke.

1.4 RStudio a základy práce v tomto prostredí

Pracovné prostredie editora **RStudio** môžeme rozdeliť do štyroch hlavných okien, ktoré uvádzame na nasledujúcom obrázku.



Obrázok 1.1: Pracovné prostredie editora RStudio

- **Okno konzoly**, nazývané aj príkazové okno, sa nachádza vľavo dole. Je ľahko odlišiteľné od ostatných, nakoľko všetky príkazy v tomto okne začínajú symbolom **>**.

Jeho použitie si možno jednoducho vyskúšať a predstaviť si ho ako akúsi kalkulačku, do ktorej je možné priamo vpisovať príkazy a po stlačení klávesu **ENTER** sa príkaz vykoná. V tomto okne je možné sa vrátiť k predchádzajúcim príkazom prostredníctvom klasických šípok na klávesnici v smere hore a dole (dozadu a dopredu).

- Okno vľavo hore predstavuje okno, v ktorom sa nachádza tzv. **skript** (zdrojový kód). Príkazy v tomto okne je možné editovať a ukladať, a preto je práca s ním oveľa praktickejšia ako používanie konzoly samostatne. Pri prvom spustení **RStudio** sa automaticky zobrazí iba okno konzoly, a je preto potrebné vytvoriť nový skript. Jeho vytvorenie iniciujeme prostredníctvom sekvencie príkazov **File** → **New File** → **R Script**, prípadne klávesovou skratkou **CTRL+SHIFT+N**². Práca v skripte pripomína poznámkový blok (notepad), pri stlačení klávesu **ENTER** sa príkaz nevykoná, ale posunieme sa na nový riadok. Ak chceme vykonať príkaz v konkrétnom riadku skriptu, je potrebné si buď celý riadok označiť, alebo sa aspoň kurzorom nachádzať v konkrétnom riadku a buď využiť tlačidlo **Run** nachádzajúce sa vpravo hore v okne skriptu, alebo klávesovú skratku **CTRL+ENTER**. Viac informácií o práci so skriptami uvádzame v podkapitole 1.10.
- Okno vpravo hore obsahuje tri záložky, konkrétne ide o **pracovné prostredie** (**environment**), **históriu** (**history**) a **prepojenia** (**connections**). Vo východiskovej pozícii sa zobrazuje pracovné prostredie, v ktorom sa nachádzajú všetky doteraz vytvorené objekty uložené v pamäti. Pri každom objekte sa nachádza dátový typ, ktorý predstavuje. Kliknutím na niektoré objekty je možné ich zobraziť (**view**) a niektoré typy dát aj priamo importovať. V histórii môžeme vidieť kompletnú históriu príkazov, ktoré boli vykonané v konzole.
- Okno vpravo dole je univerzálne a využíva sa pri práci so štandardnou grafikou a grafickými výstupmi, ktoré sa zobrazujú na karte **Plots**, na karte **Packages** je možná zjednodušená správa balíkov (inštalácia a spustenie: viac v podkapitole 1.8), karta (záložka) **Files** zobrazuje súbory a adresáre, ktoré sú v aktuálnom projekte alebo adresári a umožňuje s nimi vykonávať základné operácie (kopírovanie, vymazanie, presun...), pomocou karty **Help** zobrazujeme pomocníka, či už pre konkrétnu funkciu alebo pre celý balík funkcií a podobne (viac v podkapitole 1.6).

Pre lepšiu orientáciu v pracovnom prostredí **RStudio** odporúčame využiť aj takzvané ťaháky: **RStudio Cheat Sheets** ([odkaz tu](#)), ktoré sa týkajú rôznych oblastí, okrem základnej práce v tomto programe existujú ťaháky na množstvo oblastí, ako napríklad vizualizácia, časové rady, import údajov a podobne.

²Využívanie klávesových skratiek je veľmi vhodné pre uľahčenie práce s editorom **RStudio**. Pre kompletný zoznam skratiek je možné použiť klávesovú skratku **ALT+SHIFT+K**, prípadne je tento zoznam uvedený napríklad na stránke podpory pre **RStudio** ([odkaz tu](#)).

1.5 R ako kalkulačka – práca v konzole

V úvode si ukážeme, ako v prostredí **RStudio** pracovať prostredníctvom konzoly. V predchádzajúcej kapitole sme popisovali, kde sa okno konzoly nachádza a uviedli sme si, že všetky príkazy v tomto okne začínajú symbolom `>`. Práca v konzole je interaktívna, nakoľko všetky príkazy, ktoré do tohto okna napíšeme, sa po spustení okamžite vypíšu do konzoly, preto možno konzolu vnímať ako akúsi vysokosofistikovanú kalkulačku. Pre spustenie aktuálneho riadka príkazov, ktorý máme v konzole, používame klávesu **ENTER**. Môžeme si to vyskúšať na nasledujúcich príkladoch.

Príklad 1

Vypočítajte v **R**: $2 + 3^4 + \sin(90^\circ) + \sqrt{17} =$

Riešenie 1

Cieľom tohto príkladu je vypočítať súčet štyroch čísel. Prvé je číslo 2, pri druhom je potrebné vypočítať hodnotu mocninnej funkcie, pri ktorej číslo tri umocňujeme na štvrtú, pri treťom je potrebné určiť hodnotu goniometrickej funkcie sínus pre konkrétny uhol udaný v stupňoch a posledné číslo predstavuje druhú odmocninu z konkrétneho čísla. Pri práci s mocninnými funkciami v **R** využívame symbol striešky a pri práci s goniometrickými funkciami je ako východisková nastavená oblúčková, a nie stupňová miera, preto je potrebné 90° zmeniť na $\frac{\pi}{2}$ a využiť funkciu `sin()`. Ak chceme vypočítať druhú odmocninu, využívame na to funkciu `sqrt()`. Pre výpočet odmocnín vo všeobecnosti môžeme využívať vlastnosti mocninnej funkcie: $\sqrt[b]{x^a} = x^{\frac{a}{b}}$

```
1 > 2 + 3 ^ 4 + sin(pi / 2) + sqrt(17)
2 [1] 88.12311
```

Príklad 2

Vypočítajte v **R**: $\sin(180^\circ) - \log_3 81 \times \sqrt[4]{16^3} - \ln(e^5) =$

Riešenie 2

V tomto príklade sa stretávame s logaritmickou funkciou, na ktorú v prostredí **R** využívame funkciu `log()`, ktorá má dva argumenty, prvým je hodnota, z ktorej chceme vypočítať logaritmus (`x`) a druhým je základ logaritmu (`base`). Eulerovo číslo ($e = 2,7182818284\dots$) ako základ prirodzeného logaritmu získavame prostredníctvom funkcie `exp(1)`, ktorej jediným argumentom je exponent, na ktorý sa má Eulerovo číslo umocniť (v prípade, ak chceme samotné Eulerovo číslo, umocňujeme na prvú).

```

1 > sin(pi) - log(x = 81, base = 3) * 16 ^ (3 / 4) - log(x = exp(5),
    base = exp(1))
2 [1] -37

```

V tejto časti si môžeme uviesť poznámku o prednastavenom (východiskovom, defaultnom) nastavení niektorých funkcií, čo môžeme ilustrovať na príklade logaritmickej funkcie. Funkcia `log()` má ako východiskové nastavenie základ Eulerovo číslo, čo znamená, že táto funkcia bez toho, aby sme zadali argument `base`, vypočíta prirodzený logaritmus zvoleného čísla. Ak by sme chceli vypočítať dekadický logaritmus zvoleného čísla (logaritmus pri základe 10), musíme buď zmeniť argument `base`, alebo využiť špeciálnu funkciu, ktorá je na to určená: `log10()`, ktorej prednastavený základ je číslo 10.

```

1 > log(100)
2 [1] 4.60517
3 > log(100, base = 10)
4 [1] 2
5 > log10(100)
6 [1] 2

```

Z uvedeného vyplýva, že sme pri výpočte $\ln(e^5)$ mohli vychádzať z východiskového nastavenia funkcie `log()` a nemuseli sme definovať argument `base`.

```

1 > log(x = exp(5), base = exp(1))
2 [1] 5
3 > log(x = exp(5))
4 [1] 5

```

V okne konzoly je možné sa vrátiť k predchádzajúcim príkazom prostredníctvom klasických šípok na klávesnici v smere hore a dole (dozadu a dopredu). Po vykonaní zvoleného výpočtu sa na začiatku riadka konzoly zobrazí symbol `>`. Niekedy sa však môže stať, že namiesto daného symbolu sa na začiatku zobrazí symbol `+`. Znamená to, že ak pracujeme so zátvorkami, zabudli sme v príkaze na uzatváraciu zátvorku, alebo je program zamestnaný nejakým zložitým výpočtom. Ak v tomto prípade chceme zastaviť program a obnoviť konzolu do pôvodného stavu, je potrebné stlačiť klávesu `ESC`.

```

1 > 5 + 8 * (4 - 15 * (3 - 2))
2 [1] -83

```

```

1 > 5 + 8 * (4 - 15 * (3 - 2)
2 +

```

1.6 Pomoc a dokumentácia

Vzhľadom na to, že nie je možné poznať naspamäť všetky funkcie a ich argumenty (v ČJ parametre), má prostredie **R** veľmi dobre prepracovanú nápovedu a dokumentáciu k jednotlivým funkciám a celým balíkom. Výhodou je, že táto nápoveda je dostupná bez nutnosti pripojenia na internet.

1.6.1 Dokumentácia k funkciám

Ak chceme poznať nápovedu ku konkrétnej funkcii, využívame na to funkciu `help()`. Argumentom tejto funkcie je meno funkcie, o ktorej sa chceme dozvedieť informácie v úvodzovkách. Prípadne môžeme použiť symbol otáznika (`?`), za ktorým nasleduje meno funkcie bez úvodzoviek. V prostredí **RStudio** sa k dokumentácii ku konkrétnej funkcii dostaneme aj stlačením klávesy **F1** po tom, ako napíšeme názov funkcie. Ďalšou možnosťou, ktorú ponúka editor **RStudio**, je použitie samostatného okna nápovedy, ktoré sa nachádza vpravo dole, ktoré sme nazvali ako univerzálne okno, nakoľko okrem nápovedy toto okno ponúka aj prácu s grafickými výstupmi, balíkmi a podobne. Práca s nápovedou v tomto okne je veľmi intuitívna, je možné vyhľadávať informácie o konkrétnej funkcii rovnako ako aj vyhľadávať nejaké slovné spojenie v konkrétnej téme (**Find in Topic**). História vyhľadávania sa ukladá, je preto možné kedykoľvek sa jednoducho vrátiť k informáciám o funkciách, ktoré sme hľadali v minulosti.

Uvažujme, že chceme poznať informácie o funkcii `log()`.

```
1 help("log")
```

```
1 ?log
```

Dokumentácia, ktorá sa zobrazí, má jednotnú štandardizovanú štruktúru:

- V hlavičke dokumentácie vľavo je uvedený názov funkcie a za ním v kučeravých (zložených: `{}`) zátvorkách balík, ktorý danú funkciu obsahuje. V prípade nami uvažovanej funkcie `log()` je tam uvedené `log {base}`, čo znamená, že funkcia je súčasťou balíka **base**.
- Ďalej nasleduje názov funkcie, v našom prípade **Logarithms and Exponentials**.
- V časti **Description** je uvedený stručný popis toho, na čo je funkcia určená.
- Časť **Usage** predstavuje konkrétnu syntax pre použitie danej funkcie, aké môže mať argumenty a aké sú východiskové hodnoty týchto argumentov, pokiaľ sú nadefinované. V našom prípade je východiskovým argumentom **base** Eulerovo číslo ako základ prirodzeného logaritmu.

- V časti **Arguments** sú popísané agrumenty (parametre) funkcie ešte podrobnejšie. Okrem toho, aké sú východiskové hodnoty týchto argumentov, je tu aj uvedené akého dátového typu musia byť tieto argumenty. V našom prípade môže ísť o numerické alebo komplexné hodnoty, respektíve vektory numerických alebo komplexných hodnôt.
- Niektoré funkcie v svojej dokumentácii nemajú uvedenú časť s názvom **Details** (pozri napríklad funkciu `mean()`). Ak je to potrebné tak v tejto časti sú uvedené ešte nejaké dodatočné informácie o konkrétnej funkcii súvisiace napríklad s kompatibilitou a podobne.
- Časť **Value** vysvetľuje, aké hodnoty funkcia vracia na výstupe a akého sú dátového typu. Napríklad v našom prípade je tam uvedené, že logaritmus nuly na výstupe vracia mínus nekonečno a že logaritmus zo záporného čísla nie je definovaný (vracia na výstupe NaN - not a number).

```
1 > log(0)
2 [1] -Inf
```

```
1 > log(-1)
2 [1] NaN
3 Warning message:
4 In log(-1) : NaNs produced
```

- **S4 methods** je časť, ktorá je uvedená iba pri funkciách, ktorých sa týka objektovo orientovaný systém S4. **R** má tri objektovo orientované systémy: S3, S4 a R5.
- V časti **Source** je uvedený zdroj algoritmu výpočtu danej funkcie. V prípade jednoduchých funkcií (napríklad `mean()`) táto časť nie je uvedená.
- Časť **References** je venovaná odkazom na literatúru, ktorá pojednáva o nami uvažovanej funkcii.
- V časti **See Also** sú uvedené odkazy na iné funkcie, údaje alebo balíky, ktoré súvisia s nami uvažovanou funkciou. V našom prípade ide o odkaz na trigonometrické funkcie, ďalšie matematické funkcie (odmocnina a absolútna hodnota) a aritmetické operátory.
- Na konci dokumentácie sú uvedené príklady (**Examples**) použitia konkrétnej funkcie, ktoré sú veľmi vhodné v prípade, ak sa chceme naučiť s danou funkciou pracovať. Ak nás zaujímajú iba príklady, môžeme na to využiť samostatnú funkciu `example()`, ktorej argumentom je názov funkcie v úvodzovkách. Ukážeme si to na príklade funkcie `log()`, jednotlivé funkcie, ktoré tento príklad generuje (napríklad funkcia `cbind()`, sú vysvetlené v ďalších častiach tohto textu.

```

1 > example("log")
2
3 log> log(exp(3))
4 [1] 3
5
6 log> log10(1e7) # = 7
7 [1] 7
8
9 log> x <- 10^-(1+2*1:9)
10
11 log> cbind(x, log(1+x), log1p(x), exp(x)-1, expm1(x))
12
13      x
14 [1,] 1e-03 9.995003e-04 9.995003e-04 1.000500e-03 1.000500e-03
15 [2,] 1e-05 9.999950e-06 9.999950e-06 1.000005e-05 1.000005e-05
16 [3,] 1e-07 1.000000e-07 1.000000e-07 1.000000e-07 1.000000e-07
17 [4,] 1e-09 1.000000e-09 1.000000e-09 1.000000e-09 1.000000e-09
18 [5,] 1e-11 1.000000e-11 1.000000e-11 1.000000e-11 1.000000e-11
19 [6,] 1e-13 9.992007e-14 1.000000e-13 9.992007e-14 1.000000e-13
20 [7,] 1e-15 1.110223e-15 1.000000e-15 1.110223e-15 1.000000e-15
21 [8,] 1e-17 0.000000e+00 1.000000e-17 0.000000e+00 1.000000e-17
22 [9,] 1e-19 0.000000e+00 1.000000e-19 0.000000e+00 1.000000e-19

```

- Dokumentácia niektorých funkcií môže obsahovať aj ďalšie časti napríklad **Implementation limits**: limity použitia alebo **Note**: poznámky a podobne, v ktorých sú ešte dovysvetlené ďalšie detaily a obmedzenia použitia funkcie.
- Často sa na konci dokumentácie najmä pri novších funkciách uvádza v časti **Author(s)** aj emailový kontakt na autorov danej funkcie.

Poznámka: Vo vyššie uvedenom výstupe z programu **R** sú uvedené čísla v takzvanom vedeckom tvare (zápise). Ide napríklad o čísla **1e7** alebo **1e-03**. Ide o častý spôsob reprezentácie čísel na kalkulačkách a počítačoch. Ide o zápis v tvare **aeb**, kde **a** je desatinné číslo nazývané mantisa, **b** je celé číslo nazývané charakteristika čísla, prípadne exponent. Mantisa je číslo v rozmedzí od 1 (vrátane) do 10 (okrem), exponent rozsahom obmedzený nie je. Zápis **aeb** predstavuje:

$$a \times 10^b$$

Preto v našom prípade platí, že $1e7 = 1 \times 10^7 = 10\,000\,000$ a $1e-03 = 1 \times 10^{-3} = 0,001$.

Ak nechceme, aby program **R** zobrazoval čísla vo vedeckom zápise, môžeme na to použiť funkciu **format()** s argumentom **scientific = FALSE** alebo ešte lepšie riešenie je hneď na začiatok výpočtu nastaviť **options(scipen = 999)**, čím sa vedecký zápis čísel vo všetkých nasledujúcich príkazoch eliminuje:


```

1 > 16000000000000000
2 [1] 1.6e+15
3 > format(16000000000000000, scientific = FALSE)
4 [1] "16000000000000000"
5 > options(scipen = 999)
6 > 16000000000000000
7 [1] 16000000000000000

```

1.6.2 Dokumentácia k balíkom

Ak chceme poznať dokumentáciu k celému balíku spolu so zoznamom všetkých funkcií, ktoré obsahuje, môžeme opäť použiť funkciu `help()` s argumentom `package`, pričom v úvodzovkách je meno balíka. Dokumentáciu k balíkom je možné v prostredí RStudio získať tiež prostredníctvom univerzálneho okna vpravo dole, kde na karte **Packages** je potrebné kliknúť na meno balíka.

Uvažujme, že chceme poznať informácie o balíku `base`.

```

1 > help(package = "base")

```

1.6.3 Demonštračné kódy

Niektoré knižnice obsahujú aj takzvané demonštračné kódy, ktorými možno interaktívne zobraziť určitú funkcionálnosť. Ak chceme poznať zoznam demonštračných kódov konkrétnej knižnice, používame funkciu `demo()` s argumentom `package`, pričom v úvodzovkách je meno balíku.

Uvažujme, že chceme poznať zoznam demonštračných kódov balíku `base`.

```

1 > demo(package = "base")

```

Vidíme, že balík `base` ponúka nasledujúce štyri demonštračné kódy: `error.catching`, `is.things`, `recursion` a `scoping`. Ak si chceme spustiť niektorý z nich, opäť použijeme funkciu `demo()` a ako argument uvedieme konkrétny demonštračný kód v úvodzovkách, na čo nás môže vyzvať konzola k inicializácii daného kódu prostredníctvom stlačenia klávesu `ENTER`.

```

1 > demo("scoping")
2
3
4 demo(scoping)
5 ---- ~~~~~
6

```

```
7 Type <Return> to start :
```

1.6.4 Dokumentácia ku konkrétnej téme

Často potrebujeme poznať informácie o nejakej téme vo všeobecnosti, v tom prípade môžeme použiť funkciu `help.search()`, ktorá prehľadáva dokumentácie všetkých balíkov, ktoré sú nainštalované. Argumentom tejto funkcie je textový reťazec: téma, ktorá nás zaujíma v úvodzovkách. Rovnaký výsledok dosiahneme, ak použijeme symbol dvojitého otáznika ("??").

```
1 > help.search("linear model")
2 > ?? "linear model"
```

Ak máme k dispozícii internet, na získanie informácií o konkrétnej téme môžeme použiť funkciu `RSiteSearch()`. Argumentom tejto funkcie je opäť textový reťazec: téma, ktorá nás zaujíma v úvodzovkách. Ak uvedieme viacero kľúčových slov, vyhľadávať bude zhodu s ktorýmkoľvek z týchto kľúčových slov. Ak chceme vyhľadávať zhodu s celým slovným spojením, môžeme to doceliť tým, že toto slovné spojenie dáme do kučeravých zátvoriek. Výsledky vyhľadávania sa zobrazia v internetovom prehliadači.

```
1 > RSiteSearch("linear model")
2 A search query has been submitted to http://search.r-project.org
3 The results page should open in your browser shortly
4 > RSiteSearch("{linear model}")
5 A search query has been submitted to http://search.r-project.org
6 The results page should open in your browser shortly
```

1.6.5 Vinety

Špeciálnym typom nápovedy sú takzvané vinety. Ide o dokumentáciu, ktorá sa nevenuje špecifickej funkcii, ale ukazuje, ako pracovať s konkrétnou knižnicou ako celkom. Zoznam vinet prítomných v danom balíku získavame prostredníctvom funkcie `vignette()` s argumentom `package`, pričom v úvodzovkách je meno balíka. Zoznam týchto vinet je tiež zobrazený v dokumentácii konkrétneho balíka.

Uvažujme, že chceme poznať zoznam vinet balíka `base`.

```
1 > vignette(package = "base")
2 no vignettes found
```

Ako môžeme vidieť, balík `base` neobsahuje žiadne vinety.

Vignety obsahujú najmä novšie balíky, ktoré je potrebné stiahnuť, nainštalovať a spustiť (pre viac informácií o balíkoch pozri Kapitolu 1.8). Jedným z nich, ktorý obsahuje vinety,

je balík `survival`. Pre zobrazenie vinety používame syntax s názvom vinety v úvodzovkách a názvom konkrétneho balíku.

Uvažujme, že chceme spustiť vinetu `timedep` balíku `survival`.

```
1 > vignette(package = "survival")
2 Vignettes in package survival:
3
4 adjcurve           Adjusted Survival Curves (source, pdf)
5 concordance        Concordance (source, pdf)
6 tests              Cox models and ‘‘type 3’’ Tests (source,
7                    pdf)
8 multi              Multi-state example (source, pdf)
9 compete            Multi-state models and competing risks
10                   (source, pdf)
11 population         Population contrasts (source, pdf)
12 tiedtimes          Roundoff error and tied times (source, pdf)
13 splines            Splines, plots, and interactions (source,
14                   pdf)
15 survival           The survival package (source, pdf)
16 timedep            Using Time Dependent Covariates (source,
17                   pdf)
18 validate           Validation (source, pdf)
19 > vignette("timedep", package = "survival")
```

1.6.6 Dokumentácia z externých zdrojov

Samozrejme, že informácie je možné vyhľadať aj z externých zdrojov bez nutnosti použitia **R** a funkcií na to určených. Ideálnym zdrojom týchto informácií je internet. Pri klasickom používaní internetového vyhľadávača vo forme "how to ... in R" väčšinou prvé a najviac relevantné odkazy smerujú na stránku Stack Overflow ([odkaz tu](#)), ktorá predstavuje online komunitu používateľov a vývojárov nielen tohto programovacieho jazyka. Nachádza sa tu množstvo otázok, ktoré sú zoradené do množstva skupín prostredníctvom tzv. tagov (na začiatku roka 2020 bolo tagom "r" označených vyše 300 000 otázok). Ku každej otázke nasleduje diskusia, kde sa môže vyjadriť každý, čím môže pomôcť váš kód vylepšiť a optimalizovať. Vzhľadom na obrovské množstvo otázok a problémov, ktoré sa v rámci komunity riešia, je málokedy potrebné zakladať novú otázku, väčšinou stačí vyhľadať už predtým diskutovanú problematiku. Tento zdroj informácií využíva väčšina používateľov **R**.

Pre prácu v **R** a online vyhľadávanie na stránkach relevantných pre prácu v tomto jazyku je veľmi vhodná stránka rseek ([odkaz tu](#)), kde môžete nájsť na jednom mieste informácie z viacerých zdrojov, pričom si môžete byť istí, že sa týkajú práve tohto programovacieho jazyka.

Ďalšou online možnosťou na prístup k informáciám je napríklad stránka dokumentácie jednotlivých balíkov ([odkaz tu](#)) alebo súhrnná stránka Help for R ([odkaz tu](#)), z ktorej možno pristupovať k rôznym zdrojom týkajúcich sa **R**.

1.7 Najčastejšie chyby pri práci s R

Pri práci v **R** sa nevyhneme chybám, na ktoré nás však program stále upozorní prostredníctvom vypísania oznámenia o chybe v konzole. Okrem toho nás na chybu červeným podfarbením upozorní aj samotný editor **RStudio** ešte pred spustením samotného príkazu, v ktorom je chyba. V tejto kapitole preto uvádzame zoznam najčastejších chybových hlášok spolu s popisom, čo znamenajú a návodom, ako ich odstrániť. Pri každej chybe odkazujeme aj na konkrétne kapitoly tejto knihy, ktoré sa tej ktorej problematike venujú podrobnejšie.

- `No such file or directory` alebo `Cannot change working directory`
 - Táto chybová hláška sa vypisuje v prípade, ak je nejaký problém s pracovným adresárom. Je potrebné skontrolovať, či adresa pracovného adresára je v poriadku a či súhlasí názov súboru. Viac informácií nájdete v Kapitolách 1.9 a 7.1.
- `Object 'x' not found`
 - Objekt `x` nebol zatiaľ definovaný. Definujte objekt `x` alebo použite úvodzovky, ak chcete `"x"` použiť ako textový reťazec. Viac informácií nájdete v Kapitole 2.
- `Argument 'x' is missing without default`
 - Vo funkcii nebol zadaný argument `x`, pričom tento argument nemá východiskové (defaultné) nastavenie, čo znamená, že je nevyhnutné ho pri definovaní funkcie zadať. Viac informácií nájdete v Kapitole 5.
- `+`
 - Prvou možnosťou, prečo sa v konzole vypíše tento symbol je, že **R** stále pracuje (väčšinou ak počítame nejaký náročný výpočet, napríklad algoritmus, pri ktorom pracujeme s veľkým množstvom rôznych kombinácií). Druhou možnosťou je to, že sme zabudli na uzatváracie zátvorky. Riešením prvej možnosti je počkať, kým sa výpočet dokončí. O tom, že práve prebieha nejaký výpočet, nás upozorňuje v **RStudio** červená značka s nápisom **STOP** v hornej pravej časti konzoly. Riešením druhej možnosti je napísať do konzoly uzatváraciu zátvorku. Riešením oboch prípadov je stlačenie klávesy **ESC**, čím sa výpočet zastaví. Viac informácií nájdete v Kapitole 1.5.

-
- Unexpected '}' in ")" or Unexpected '}' in "}"
 - V tomto prípade ide o opačnú situáciu v porovnaní s predchádzajúcou chybovou hláškou (+). Snažíme sa uzatvoriť niečo, čo ešte nebolo otvorené. Riešením je pridanie otváracích zátvoriek. Viac informácií nájdete v Kapitole 1.5.
 - Unexpected 'else' in "else"
 - S touto chybou sa stretávame v prípade, ak pracujeme s podmienkami typu `if()`. Riešením je napísanie `else`, ktoré sa viaže na podmienku typu `if` na rovnaký riadok, v ktorom sa nachádza posledná zátvorka definujúca, čo sa má urobiť, ak je podmienka splnená: `} else {`. Viac informácií nájdete v pokračovaní tejto knihy (*R snadno a rychle 2 Vizualizace dat a programování*).
 - Missing value where TRUE/FALSE needed
 - Opäť chyba týkajúca sa podmienky typu `if()`. Táto situácia nastáva vo chvíli, keď podmienka, ktorá má nasledovať v klasických zátvorkách za `if`, bola buď chybné definovaná, alebo chýba úplne. Viac informácií nájdete v pokračovaní tejto knihy (*R snadno a rychle 2 Vizualizace dat a programování*).
 - The condition has length > 1 and only the first element will be used
 - Opäť chyba týkajúca sa podmienky typu `if()`. Táto chyba nastáva vtedy, keď do klasických zátvoriek za `if` vložíme atomický vektor logických premenných namiesto jednej konkrétnej logickej hodnoty (`TRUE` / `FALSE`). Viac informácií nájdete v pokračovaní tejto knihy (*R snadno a rychle 2 Vizualizace dat a programování*).
 - Non-numeric argument to binary operator
 - S touto chybou sa stretávame, ak sa snažíme robiť matematické výpočty s nejakým objektom, ktorý nie je číslo (nie je typu `numeric`). Na transformáciu na tento dátový typ môžeme použiť funkciu `as.numeric()`. Viac informácií nájdete v Kapitole 3.
 - Argument is of length zero alebo Replacement is of length zero
 - Objekt, s ktorým chceme pracovať, je prázdny (`NULL`), napríklad prázdny vektor vytvorený iba pomocou funkcie `c()`, bez definovania hodnôt. Riešením je zdefinovanie objektu nanovo. Viac informácií nájdete v Kapitole 6.

1.8 Práca s balíkmi

Pri praktických analýzach v **R** často potrebujeme zvýšiť funkčnosť tohto nástroja prostredníctvom tzv. balíkov (knížnic). Po základnej inštalácii programu **R** a **RStudio**

sa po zapnutí programu automaticky spustia (načítajú) iba niektoré balíky. Ide o balíky ako `base`, `datasets`, `graphics`, `methods`, `stats` a podobne, ktoré nám umožňujú vykonať základné kalkulácie, ale možnosť použitia sofistikovanejších metód prostredníctvom týchto základných balíkov nie je možná.

Balík možno chápať ako nejakú sadu funkcií, ktorá sa väčšinou viaže na konkrétnu oblasť. Napríklad balík `ggplot2` predstavuje sadu funkcií na vizualizáciu údajov, balík `tidyquant` obsahuje funkcie na kvantitatívnu analýzu údajov z finančných trhov a podobne.

Princíp práce s balíkmi je veľmi jednoduchý. Pokiaľ sme ešte s konkrétnym balíkom nepracovali a nemáme ho nainštalovaný, je potrebné ho nainštalovať. Inštaláciu postačí vykonať iba raz. Samotnou inštaláciou však ešte balík nie je funkčný. Preto je potrebné pred aplikáciou funkcií z konkrétného balíka predtým balík načítať, respektíve spustiť pri každom spustení **R**.

Inštalácia balíka v prostredí **RStudio** je veľmi jednoduchá. Prvou možnosťou je využitie univerzálneho okna vpravo dole, kde v časti **Packages** klikneme na tlačidlo **Install** a do okna vpíšeme názov balíka, ktorý chceme nainštalovať. Inou možnosťou je priamo funkcia na inštalovanie `install.packages()`, ktorej hlavným argumentom je názov balíka v úvodzovkách. Medzi ďalšie argumenty patrí napríklad to, či chceme nainštalovať aj balíky, z ktorých balík vychádza (**dependencies**) alebo zložku, do ktorej sa bude balík inštalovať pomocou argumentu `lib = " "`, kde v úvodzovkách je uvedená adresa zložky, do ktorej chceme balík inštalovať. Ak chceme napríklad nainštalovať balík `ggplot2` spolu s ostatnými balíkmi, na ktoré je naviazaný:

```
1 > install.packages("ggplot2", dependencies = TRUE)
```

Inštaláciu je potrebné urobiť len raz, niekedy je však potrebné balíky aktualizovať a pracovať s aktuálnou verziou balíka, ktorá obsahuje najnovšie funkcie a podobne. Na to slúži funkcia `update.packages()`, jednoduchší spôsob aktualizácie balíkov je využiť univerzálne okno vpravo dole, kde v časti **Packages** klikneme na tlačidlo **Update**. Zobrazí sa nám tabuľka nainštalovaných balíkov, ktoré možno aktualizovať spolu s aktuálnou verziou a najnovšou verziou balíka. Jednoduchým zaškrtnutím jednotlivých balíkov a kliknutím na tlačidlo **Install Updates** je možné aktualizáciu konkrétného balíka vykonať.

Spustenie (načítanie) balíka je potrebné stále pred prvou aplikáciou nejakej funkcie, ktorá sa v balíku nachádza. Je to z dôvodu šetrenia pamäte, nakoľko nie každý nainštalovaný balík je potrebné mať spustený pri realizácii konkrétnej analýzy. Prvou možnosťou na spustenie balíka je využitie univerzálneho okna vpravo dole, kde v časti **Packages** máme uvedený zoznam všetkých nainštalovaných balíkov, spolu s ich popisom (**Description**) a verziou (**Version**). Vpravo sú dve tlačidlá, prvým z nich v tvare **+** je možné zobrazíť si online dokumentáciu konkrétného balíka, druhým z nich v tvare **X** je možné balík odinštalovať, čo však neodporúčame, hlavne pokiaľ ide o vstavané balíky, prípadne balíky, z ktorých vychádzajú iné balíky, čím by sa mohla znemožniť ich funkcionálnosť. Rovnaký efekt dosiahneme, ak využijeme funkciu `remove.packages()`, ktorej hlavným argumen-

tom je názov balíka v úvodzovkách.

Ak chceme spustiť konkrétny balíček, je potrebné označiť zaškrtávací box, ktorý sa nachádza pred názvom balíka. Inou možnosťou je priamo funkcia na spustenie `library()`, ktorej hlavným argumentom je názov balíka bez úvodzoviek. Pri inštalácii alebo spustení balíkov môžu v konzole vyskakovať nejaké červené varovné nápisy, ktoré rozdeľujeme do dvoch skupín:

- Error: chyba, zadaný kód nie je možné realizovať.
- Warning message: varovné oznámenie, kód je možno vykonať, program nás však informuje o niečom, čo sme mohli prehliadnuť a na čo by sme si mali dať pozor.

Ak chceme napríklad spustiť už nainštalovaný balík `ggplot2`:

```
1 > library(ggplot2)
2 Warning message:
3 package ggplot2 was built under R version 3.6.2
```

V tomto prípade nás varovný odkaz upozorňuje na to, v akej verzii **R** bol balík vytvorený, čím nás chce upozorniť na to, že ak nemáme danú verziu **R**, môžeme mať problémy s funkcionalitou a niektoré funkcie balíka nemusia správne fungovať.

Ak chceme napríklad spustiť predtým nenainštalovaný balík `tidyquant`:

```
1 > library(tidyquant)
2 Error in library(tidyquant) : there is no package called tidyquant
```

V tomto prípade ide o chybu, vzhľadom na to, že daný balík nie je nainštalovaný, nie je možné ho spustiť.

Zoznam nainštalovaných balíkov je možné vidieť v univerzálnom okne vpravo dole v časti **Packages**, pričom spustené (načítané) balíky majú zapnuté začiarkovacie políčko (zaškrtávatko). Okrem toho je možné zoznam nainštalovaných balíkov zobrazíť aj prostredníctvom funkcie `installed.packages()`.

Ak potrebujeme spúšťať nejakú sekvenciu kódu na viacerých zariadeniach, pričom na niektorých z nich je nainštalovaný nejaký konkrétny balík potrebný k analýze a na niektorých nie, odporúčame nasledovnú syntax, ktorá na začiatku overí, či je potrebné daný balík inštalovať a ak áno, nainštaluje ho a následne spustí, ak nie, tak ho iba spustí a proces inštalácie sa preskočí, nakoľko nie je potrebný. Ukážeme si to na príklade balíka `tidyquant`:

```
1 > if (!require("tidyquant")) {
2 +   install.packages("tidyquant", dependencies = TRUE)
3 +   library(tidyquant)
4 + }
```

Balíky v **R** sú k dispozícii v centralizovaných repozitároch. Poznáme tieto tri hlavné centrálné repozitáre balíkov:

- CRAN ([odkaz tu](#)): ide o najväčší a hlavný centrálny repozitár balíkov, s ktorými pracujeme. Balíky, ktoré sa tu nachádzajú, prešli kontrolou komunity používateľov, sú už väčšinou vo finálnej fáze. Tento repozitár je aj prednastavený na sťahovanie balíkov v prostredí **RStudio**.
- GitHub ([odkaz tu](#)): tento repozitár obsahuje nové balíky, ktoré sa zatiaľ nedostali na CRAN a sú väčšinou ešte vo vývojovej fáze.
- Bioconductor ([odkaz tu](#)): tento repozitár obsahuje balíky určené na analýzu genomických údajov.

Na konci tejto časti ešte upozorňujeme na to, že pri aktualizácii **R** na nejakú novšiu verziu je potrebné balíky nainštalovať nanovo.

1.9 Vytvorenie pracovného adresára

Ešte predtým, ako si vysvetlíme ako pracujeme so skriptami, je vhodné nastaviť si pracovný adresár (priečinok, v ČJ složka). Ide o priečinok v počítači, v ktorom aktuálne pracujeme. Znamená to, že pri importe (načítaní) súborov do **R** sa načítavané súbory nachádzajú v tomto priečinku a pri exporte (ukladaní) súborov sa súbory primárne ukládajú do tohto priečinka (môže ísť napríklad o datasety, ale aj o históriu príkazov alebo grafické výstupy).

Pri prvom spustení **RStudio** je väčšinou pracovný adresár nastavený ako priečinok Dokumenty (Documents), prípadne iný všeobecný priečinok v závislosti od používaného operačného systému. Ak otvoríme **RStudio** spustením konkrétneho skriptu, ako pracovný adresár sa nastaví priečinok, v ktorom sa daný skript nachádza. Ak teda napríklad otvoríme skript, ktorý máme uložený na pracovnej ploche, automaticky sa pracovná plocha nastaví ako pracovný priečinok. Získanie aktuálneho pracovného priečinka môžeme získať prostredníctvom funkcie `getwd()`.

```
1 > getwd()
2 [1] "D:/Users/..."
```

Ak chceme nastaviť pracovný priečinok, využívame na to funkciu `setwd()`, ktorej argumentom je cesta k danému priečinku. Je však potrebné upozorniť na opačný tvar lomiek pri definovaní adresy ako v prípade adresy v OS Windows (Windows: "\", **R**: "/").

```
1 > setwd("D:/Users")
```


Ak cesta k pracovnému priečinku je dlhšia, tak s cieľom eliminácie prípadných chýb vplyvom zlej adresy a rôznych operačných systémov odporúčame používať funkciu `file.path()`, ktorá spojí jednotlivé časti cesty k priečinku tak, aby fungovali na každom operačnom systéme:

```
1 > file.path("D:", "Users", "Programs", "R", "Files")
2 [1] "D:/Users/Programs/R/Files"
3 > setwd(file.path("D", "Users", "Programs", "R", "Files"))
```

Jednoduchším spôsobom, ako je možné v prostredí **RStudio** zistiť aktuálny pracovný priečinok, respektíve zmeniť ho na iný, je využiť univerzálne okno vpravo dole, kde v časti **Files** je možné vidieť cestu k pracovnému priečinku a aj celý obsah tohto priečinka. Adresu vieme zmeniť jednoduchým kliknutím, avšak ak chceme, aby sa zmenený priečinok stal pracovným priečinkom, musíme využiť tlačidlo **More** a kliknúť na **Set As Working Directory**. V tejto časti **Files** je okrem iného možné pracovať podobne ako v akomkoľvek správcovi súborov, je možné kopírovanie súborov, vymazávanie, presun, premenovanie, zobrazenie skrytých súborov alebo vytvorenie nového priečinka.

Asi najpoužívanejším spôsobom na nastavenie pracovného priečinka je využitie sekvencie príkazov **Session** → **Set Working Directory** → **Choose Directory**, prípadne klávesovou skratkou **CTRL+SHIFT+H**.

1.10 Vytvorenie skriptu

V Kapitole 1.5 sme si uviedli, ako pracovať v konzole a ukázali sme si, že **R** môže pracovať ako jednoduchá kalkulačka. Pri práci v **R** však oveľa častejšie ako písanie príkazov do konzoly využívame tvorbu a spúšťanie takzvaných skriptov.

Skript predstavuje jednoduchý textový súbor **R** príkazov, pričom tieto príkazy nielen z dôvodov väčšej prehľadnosti píšeme každý na nový riadok. Ide o súbor uložený vo formáte **nazevSouboru.R**, ktorý je kedykoľvek editovateľný a ktorý nám umožňuje spúšťať nami vytvorené príkazy, analýzy a podobne. Okrem toho v skripte prostredníctvom symbolu mriežky (**#**) môžeme akúkoľvek časť skriptu komentovať, čo znamená, že čokoľvek, čo sa v rámci konkrétneho riadka skriptu nachádza za týmto symbolom, bude **R** ignorovať.

Na vytvorenie nového skriptu používame sekvenciu príkazov **File** → **New File** → **R Script**, prípadne klávesovú skratku **CTRL+SHIFT+N**. Samotný skript je možné editovať napríklad aj v poznámkovom bloku (notepad), avšak **RStudio** je veľmi vhodný editor, nakoľko dokáže farebne zvýrazniť zátvorky a inú syntax pre lepšiu prehľadnosť, pomôcť pri nápovedi o názve funkcie a jej parametroch, odhaliť niektoré chyby a podobne.

K spusteniu celého skriptu, ktorý sa nachádza v pracovnom priečinku (nemusí ísť o aktuálny skript), slúži funkcia `source()`, ktorej jedným z argumentov je názov skriptu, pričom je možné nastaviť aj cestu k tomuto skriptu, čo znamená, že sa skript nemusí nachádzať

v aktuálnom pracovnom priečinku. Všetky riadky skriptu sa spustia jeden za druhým a vypíšu sa v konzole. Ak nechceme, aby sa do konzoly vypisovali tie časti kódu, ktoré sa práve vyhodnocujú, môžeme využiť argument `echo`.

```
1 > source("navezSouboru.R") #ak ide o skript z Working Directory
```

```
1 > source("D:/Users/.../navezSouboru.R") #ak ide o skript ulozeny  
inde ako vo Working Directory
```

Ak chceme v skripte spustiť iba nejakú jeho časť, využívame na to tlačidlo **Run** v pravej hornej časti okna skriptu. Rovnaký efekt má použitie klávesovej skratky **CTRL+ENTER**. Ak sa kurzorom nachádzame na konkrétnom riadku skriptu, spustí sa tento riadok, ak máme označený väčší blok príkazov na viacerých riadkoch, spustia sa všetky tieto príkazy. Najčastejšie pri práci so skriptom používame nasledujúce klávesové skratky:

- **CTRL+ENTER**: spustí aktuálny riadok alebo označenie príkazov,
- **CTRL+SHIFT+ENTER**: spustí všetky riadky (celý aktuálny skript),
- **CTRL+ALT+B**: spustí všetky riadky skriptu od prvého po aktuálny riadok,
- **CTRL+ALT+E**: spustí všetky riadky skriptu od aktuálneho riadka po posledný riadok.

Pri písaní skriptov sa z dôvodu prehľadnosti, čitateľnosti a jednoduchšieho nájdenia chýb odporúča dodržiavať určitý štýl a skripty formálne upravovať (napríklad používať rovnaké medzery a podobne). **RStudio** umožňuje skript preformátovať prostredníctvom tlačidla **Code Tools** (symbol čarovného prútika) vľavo hore okna skriptu. Sú možné viaceré úpravy skriptu, asi najpraktickejšie je preformátovať celý kód (**Reformat Code**: klávesová skratka **CTRL+SHIFT+A**) a zarovnať komentáre (**Reflow Comment**: klávesová skratka **CTRL+SHIFT+)**).

Kapitola 2

Vytváranie objektov

R je objektovo orientovaný procedurálny programovací jazyk. Všetko, čo v tomto prostredí existuje a s čím pracujeme, je objekt. Z hľadiska syntaxe je **R** takzvaný *case sensitive* programovací jazyk, čo znamená, že objekt **x** nie je to isté, čo objekt **X**.

Objekt v **R** môže obsahovať v jednom okamihu iba jednu konkrétnu „hodnotu“, pričom môže ísť o objekt rôzneho dátového typu a rôznej dátovej štruktúry. Môže pritom ísť aj napríklad o funkciu, graf¹ a podobne. Typ objektu pri jeho inicializácii nie je potrebné nijako špeciálne definovať.

Uvažujme, že chceme vytvoriť objekty **a**, **b**, **c** tak, že do nich priradíme čísla 1, 2, 3. Hodnoty sa do objektov priraďujú nasledovným spôsobom:

- Prostredníctvom šípky (**->**), kde šípka stále ukazuje k názvu objektu a na druhej strane je „hodnota“, ktorú chceme do objektu priradiť. V tejto časti chceme upozorniť na to, že priradenie opačnou šípkou, ako je to uvedené na príklade objektu **b** (**->**) nie je z dôvodu prehľadnosti a čitateľnosti kódu vhodné.

```
1 > a <- 1 #odporucane priradenie
2 > 2 -> b #neodporucane priradenie
```

- Symbol „rovná sa“ (**=**) má vo väčšine prípadov ekvivalentné použitie ako šípka (**<-**), avšak môžu nastať situácie, keď pomocou tohto symbolu nie je možné priradenie vykonať, preto odporúčame na priradenie používať symbol šípky. Príkladom, keď nie je možné priradenie pomocou „rovná sa“ vykonať, je napríklad situácia, keď je tento symbol použitý ako argument funkcie.

```
1 > log(x = 100, base = 10)
2 [1] 2
3 > x
```

¹Napríklad pri práci a vizualizácii s balíkom **ggplot2**.

```
4 Error: object 'x' not found
5 > log(x <- 100, base = 10)
6 [1] 2
7 > x
8 [1] 100
```

V prvom prípade sme pomocou `=` iba uviedli argument funkcie `log()`, to znamená číslo, z ktorého chceme vypočítať dekadický logaritmus. Objekt `x` sa však nevytvoril. V druhom prípade sme vytvorili objekt `x`, do ktorého sme uložili číslo 100 a z tohto čísla sme následne vypočítali dekadický logaritmus. Pre zobrazenie toho, čo sa nachádza v konkrétnom objekte, môžeme vypísať názov objektu na nový riadok. V prvom prípade, keďže sa žiadny objekt nevytvoril, nám vypísalo chybu, že zadaný objekt neexistuje. V druhom prípade nám vypísalo jeho hodnotu (100).

- Treťou a asi najmenej využívanou možnosťou priradenia je využitie funkcie `assign()`, pri ktorej názov objektu, do ktorého chceme priradiť, dávame do úvodzoviek:

```
1 > assign("c", 3)
```

Ak chceme vypísať hodnotu objektu do konzoly, môžeme to urobiť buď tak, ako sme už uviedli vyššie, že napíšeme meno tohto objektu, alebo môžeme využiť aj funkciu `print()`, ktorej argumentom je názov objektu:

```
1 > a
2 [1] 1
3 > print(b)
4 [1] 2
```

Ak chceme vypísať hodnotu objektu do konzoly hneď pri procese priradenia objektu, môžeme to urobiť tak, že celý proces dáme do obyčajných zátvoriek:

```
1 > (d <- 4)
2 [1] 4
```

Operátor priradenia má dve základné úlohy. V prípade, ak objekt neexistuje, vytvorí nový objekt a uloží doň to, čo doň uložiť chceme. V prípade, ak objekt už existuje a je v ňom niečo uložené, operátorom priradenia môžeme zmeniť to, čo sa v objekte nachádza. V objekte `a` sa momentálne nachádza číslo 1, zmeňme to a priradíme tam číslo 10. Pri operátore priradenia však musíme dávať pozor na to, aby sme náhodou neurobili zbytočnú medzeru:

```
1 > print(a)
2 [1] 1
3 > a <- 10
4 > print(a)
5 [1] 10
```

V tomto prípade sme si najprv vypísali, že v objekte **a** sa nachádza číslo 1 a následne sme do tohto objektu priradili hodnotu 10, ktorú sme potom vypísali do konzoly.

```
1 > a < -10
2 [1] FALSE
```

V tomto druhom prípade sme však kvôli medzere nepoužili operátor priradenia, ale logický operátor, pri ktorom sa pýtame, či je číslo uložené v objekte **a** menšie ako -10. Keďže vieme, že číslo uložené v objekte **a** je rovné 10, daná nerovnosť neplatí, preto je výsledkom tejto operácie **FALSE** (nepravda, nerovnosť neplatí). Vo všeobecnosti v prostredí **R** na medzerách až tak nezáleží, ale ako sme si ukázali, v prípade operátora priradenia ide o možný zdroj zbytočných chýb.

S objektmi môžeme robiť rôzne kalkulácie, napríklad:

```
1 > a + b + c + d + x #10+2+3+4+100
2 [1] 119
```

```
1 > a * b - 5 #10*2-5
2 [1] 15
```

Ak chceme priradiť rovnakú „hodnotu“ do viacerých objektov, môžeme použiť nasledujúcu syntax:

```
1 > x <- y <- z <- 100
```

Ak chceme vedieť, aké objekty máme aktuálne uložené v pamäti, môžeme použiť funkciu **ls()**:

```
1 > ls()
2 [1] "a" "b" "c" "d" "x" "y" "z"
```

Vytvorené objekty spolu s „hodnotami“, ktoré obsahujú, môžeme vidieť aj v okne pracovného prostredia (**Environment**) vpravo hore.

Na pripomenutie uvádzame citlivosť na veľké a malé písmená (case sensitivity). Objekt **b** sa v pamäti nachádza a je rovný 2, objekt **B** sa v pamäti nenachádza:

```
1 > print(b)
2 [1] 2
3 > print(B)
4 Error in print(B) : object 'B' not found
```

Ak chceme nejaké objekty vymazať, využívame funkciu **rm()** (remove), ktorej argumenty sú názvy objektov, ktoré chceme vymazať.

```
1 > ls()
```

```

2 [1] "a" "b" "c" "d" "x" "y" "z"
3 > rm(y, z)
4 > ls()
5 [1] "a" "b" "c" "d" "x"

```

Ak chceme vymazať všetky objekty, ktoré máme aktuálne uložené v pracovnom prostredí, môžeme na to použiť kombináciu funkcií `rm()` a `ls()`. Docielime to aj tak, keď v okne pracovného prostredia (**Environment**) vpravo hore klikneme na tlačidlo so symbolom metly (**Clear objects from the workspace**). Následne v okne pracovného prostredia (**Environment**) vpravo hore máme uvedené, že "Environment is empty" (pracovné prostredie je prázdne), o čom svedčí aj vypísanie objektov aktuálne uložených v pamäti, ktoré na výstupe vypisuje prázdny textový reťazec.

```

1 > rm(list = ls())
2 > ls()
3 character(0)

```

2.1 Pravidlá pri vytváraní objektov

Objekty, ktoré vytvárame, môžeme pomenovať prakticky ľubovoľne podľa potreby, avšak existujú niektoré výnimky, ktoré je potrebné dodržiavať. Meno objektu sa môže skladať iba z písmen, číslíc, bodiek a podčiarkovníkov, nesmie začínať číslicom, ale musí začínať písmenom alebo bodkou, za ktorou nenasleduje číslica. Taktiež existujú rezervované „slová“, ktoré taktiež nie je možné použiť na pomenovanie objektov, ide o tieto slová:

- | | | |
|-------------------------|----------------------|------------------------------|
| • <code>if</code> | • <code>next</code> | • <code>NA</code> |
| • <code>else</code> | • <code>break</code> | • <code>NA_integer_</code> |
| • <code>repeat</code> | • <code>TRUE</code> | • <code>NA_real_</code> |
| • <code>while</code> | • <code>FALSE</code> | • <code>NA_complex_</code> |
| • <code>function</code> | • <code>NULL</code> | • <code>NA_character_</code> |
| • <code>for</code> | • <code>Inf</code> | |
| • <code>in</code> | • <code>NaN</code> | |

V prostredí **R** existujú aj niektoré konštanty, napríklad `pi`, ktorých hodnoty sú vopred v programe uložené. Vytvorenie objektu s názvom týchto konštánt je síce možné, avšak neodporúčame to robiť, nakoľko tým prepíšeme tieto pôvodne existujúce hodnoty a nebudeme s nimi vedieť ďalej pracovať.

Kapitola 3

Dátové typy

Medzi základné dátové typy, s ktorými pracujeme v **R**, zaraďujeme:

- **logical**: tento dátový typ môže nadobúdať iba logické hodnoty **TRUE** (pravda) alebo **FALSE** (nepravda). Je možné tieto hodnoty skrátiť na **T**, respektíve **F**, ale vzhľadom na to, že by sme v kóde mohli tieto písmená využiť ako názvy objektov, veľmi sa toto skracovanie neodporúča.
- **integer**: tento dátový typ môže nadobúdať iba celé čísla, ako kladné, tak aj záporné. Ak chceme s číslom pracovať ako s typom **integer**, je potrebné za číslo napísať **L**, napríklad **5L**. Ak by sme **L** nenapísali, uložilo by sa to v type **double** (reálne číslo) aj napriek tomu, že ide o celé číslo.
- **double**: dátový typ, ktorý je prednastavený na prácu s reálnymi číslami. Môže teda ísť o kladné aj záporné, celé aj desatinné, racionálne aj iracionálne reálne čísla. Presnosť výpočtu (chyba zaokrúhľovania) je pri práci s dátovým typom **double** 16 desatinných miest.
- **complex**: pri tomto dátovom type sa stretávame, ak pracujeme s komplexnými číslami. Pri týchto číslach je potrebné nadefinovať ich reálnu aj imaginárnu časť pomocou **i**: napríklad **1+4i**.
- **character**: dátový typ, s ktorým sa stretávame pri práci s textovými reťazcami. V **R** s textom pracujeme tak, že ho dávame do úvodzoviek. Všetko, čo je v úvodzovkách, sa považuje za typ **character**, aj v prípade, ak by napríklad v úvodzovkách bolo nejaké číslo.

Na ilustráciu práce s dátovými typmi si vytvoríme päť objektov s názvami **x.logical**, **x.integer**, **x.double**, **x.complex**, **x.character**, do ktorých priradíme dátové typy na základe názvu objektu:

```
1 > x.logical <- TRUE
2 > x.integer <- 5L
3 > x.double <- 3 / 4
4 > x.complex <- 1 + 4i
5 > x.character <- "ahoj"
```

Ak chceme vedieť, o aký dátový typ ide, môžeme na to použiť funkciu `typeof()`, ktorá vracia typ premennej ako textový reťazec.

Logická funkcia `is.logical()` vracia hodnotu `TRUE` v prípade, ak objekt je typu `logical`. Logická funkcia `is.integer()` vracia hodnotu `TRUE` v prípade, ak objekt je typu `integer`. Logická funkcia `is.double()` vracia hodnotu `TRUE` v prípade, ak objekt je typu `double`. Logická funkcia `is.complex()` vracia hodnotu `TRUE` v prípade, ak objekt je typu `complex`. Logická funkcia `is.character()` vracia hodnotu `TRUE` v prípade, ak objekt je typu `character`.

Pre reálne čísla vo všeobecnosti (typ `integer` a `double`) zvykneme používať funkciu `is.numeric()`, ktorá vráti hodnotu `TRUE` v prípade, ak objekt je typu `integer` alebo `double`.

```
1 > typeof(x.logical)
2 [1] "logical"
3 > is.logical(x.logical)
4 [1] TRUE
5 > is.numeric(x.logical)
6 [1] FALSE
7 >
8 > typeof(x.integer)
9 [1] "integer"
10 > is.integer(x.integer)
11 [1] TRUE
12 > is.numeric(x.integer)
13 [1] TRUE
14 >
15 > typeof(x.double)
16 [1] "double"
17 > is.double(x.double)
18 [1] TRUE
19 > is.numeric(x.double)
20 [1] TRUE
21 >
22 > typeof(x.complex)
23 [1] "complex"
24 > is.complex(x.complex)
25 [1] TRUE
26 > is.numeric(x.complex)
```



```

27 [1] FALSE
28 >
29 > typeof(x.character)
30 [1] "character"
31 > is.character(x.character)
32 [1] TRUE
33 > is.numeric(x.character)
34 [1] FALSE

```

R v sebe obsahuje aj niektoré konštanty, napríklad π . Ide o iracionálne číslo, ktoré zaraďujeme medzi reálne čísla, preto je jeho typ `double`.

```

1 > pi
2 [1] 3.141593
3 > typeof(pi)
4 [1] "double"

```

3.1 Zmena dátového typu

Dátové typy možno v niektorých prípadoch meniť. Smer zmeny je daný logicky, napríklad celé číslo uložené ako typ `integer` je možné zmeniť na typ `double` a považovať ho vo všeobecnosti za reálne. Pri opačnom smere, ak typ `double` chceme zmeniť na `integer`, v prípade, ak ide o desatinné číslo, zmenou dátového typu dostávame hodnotu tohto čísla pred desatinnou čiarkou (dolnú celú časť) bez ohľadu na zaokrúhľovanie. Nižšie uvádzame príklady zmien niektorých dátových typov. Zmenu dátového typu vykonávame prostredníctvom funkcie `as.TYP()`, kde `TYP` je názov dátového typu, napríklad `as.numeric()`. Ako môžeme vidieť, ak typ `double` chceme zmeniť na `logical`, tak každé číslo či kladné, či záporné vracia logickú hodnotu `TRUE`, jedine nula vracia logickú hodnotu `FALSE`. Ak chceme, naopak, typ `logical` zmeniť na číslo, tak `TRUE` sa mení na číslo 1 a `FALSE` sa mení na číslo 0. Pomocou funkcie `as.character()` meníme pôvodný objekt na textový reťazec, takže bez ohľadu na pôvodný dátový typ sa všetko dá do úvodzoviek. Ak reálne číslo meníme na komplexné, jeho imaginárna časť sa nastaví na nulu, pri opačnej konverzii sa berie do úvahy iba reálna časť a **R** nás upozorní prostredníctvom upozornenia `Warning message: imaginary parts discarded in coercion`, že imaginárnu časť komplexného čísla ignorujeme. Ak chceme zmeniť textový reťazec na číslo, je to možné iba v prípade, ak v úvodzovkách je nejaké číslo, ak tam je normálny text, táto konverzia nie je možná a opäť dostaneme upozornenie, tentokrát `Warning message: NA introduced by coercion`.

```

1 > as.double(x.integer)
2 [1] 5
3 >
4 > as.integer(5.9)
5 [1] 5

```

```

6 > as.integer(5.1)
7 [1] 5
8 >
9 > as.logical(0.4)
10 [1] TRUE
11 > as.logical(-0.4)
12 [1] TRUE
13 > as.logical(0)
14 [1] FALSE
15 >
16 > as.numeric(TRUE)
17 [1] 1
18 > as.numeric(FALSE)
19 [1] 0
20 >
21 > as.character(TRUE)
22 [1] "TRUE"
23 > as.character(5)
24 [1] "5"
25 >
26 > as.complex(5)
27 [1] 5+0i
28 > as.numeric(1 + 4i)
29 [1] 1
30 Warning message:
31 imaginary parts discarded in coercion
32 >
33 > as.numeric("ahoj")
34 [1] NA
35 Warning message:
36 NAs introduced by coercion
37 > as.numeric("5")
38 [1] 5
39 >
40 > as.logical("5")
41 [1] NA
42 > as.logical(as.numeric("5"))
43 [1] TRUE

```

V predchádzajúcich prípadoch sme vplyvom konverzie dostali niekedy na výstupe chýbajúcu hodnotu NA (not available). Vo všeobecnosti platí, že operácie s aspoň jednou NA hodnotou dávajú stále vo výsledku hodnotu NA, preto si treba dať na to pozor. Jedným z riešení môže byť použitie funkcie `na.omit()` uvedenej v Kapitole 7.4. Okrem toho poznáme ešte iný druh chýbajúcej hodnoty, tzv. NaN (not a number), ktorá vznikne v prípade, ak prevedieme nejaké nepovolené matematické operácie (s výnimkou delenia nenulového

čísla nulou). V prípade delenia nenulového čísla nulou je výsledok nevlastné číslo `Inf`, respektíve `-Inf`. Tieto nevlastné čísla **R** považuje za reálne (typ `double`).

```
1 > 0 / 0
2 [1] NaN
3 >
4 > 5 / 0
5 [1] Inf
6 > - 5 / 0
7 [1] -Inf
8 >
9 > typeof(Inf)
10 [1] "double"
11 >
12 > Inf / Inf
13 [1] NaN
```

Kapitola 4

Operátory

S rôznymi dátovými typmi môžeme vykonávať množstvo operácií pomocou operátorov a funkcií. Operátory rozdeľujeme do viacerých skupín, pričom uvedieme tie, ktoré v **R** využívame najčastejšie.

4.1 Aritmetické operátory

Aritmetické operátory používame pri jednoduchých matematických operáciách, ako napríklad spočítavanie, umocňovanie a podobne.

+	operátor sčítania,
-	operátor odčítania,
*	operátor násobenia,
%%	operátor maticového súčinu,
/	operátor delenia,
^	operátor umocnenia,
%/%	operátor celočíselného delenia,
%%	operátor zvyšku po delení (modulo).

Aritmetické operátory používame pri klasických matematických operáciách. Platia všetky pravidlá, ktoré poznáme z matematiky (zátvorka má prednosť, násobenie a delenie má prednosť pred sčítaním a odčítaním a podobne).

```
1 > 5 + 8 * 9
2 [1] 77
3 > 5 + 72
4 [1] 77
5 >
6 > (5 + 8) * 9
7 [1] 117
```

```

8 > 13 * 9
9 [1] 117
10 >
11 > 426 / 6 - 3
12 [1] 68
13 > 71 - 3
14 [1] 68
15 >
16 > 426 / (6 - 3)
17 [1] 142
18 > 426 / 3
19 [1] 142
20 >
21 >
22 > 2 ^ 5 - 3 * 8
23 [1] 8
24 > 32 - 24
25 [1] 8
26 >
27 > (2 ^ 5 - 3) * 8
28 [1] 232
29 > 29 * 8
30 [1] 232
31 >
32 > 9 %/ 4
33 [1] 2
34 >
35 > 9 %% 4
36 [1] 1

```

4.2 Relačné (porovnávacie) operátory

Pri relačných operátoroch porovnávame číselné hodnoty objektov. Keďže ide stále o otázky s jednoznačnou odpoveďou áno alebo nie, výstupom porovnávacích operátorov je logická hodnota `TRUE` alebo `FALSE`.

<code>a == b</code>	Je <code>a</code> rovné <code>b</code> ? (nepomýliť si to s operátorom priradenia =)
<code>a != b</code>	Je <code>a</code> nerovné <code>b</code> ? (symbol <code>!</code> používame pre negáciu akéhokoľvek výrazu)
<code>a < b</code>	Je <code>a</code> menšie ako <code>b</code> ?
<code>a > b</code>	Je <code>a</code> väčšie ako <code>b</code> ?
<code>a <= b</code>	Je <code>a</code> menšie alebo rovné ako <code>b</code> ?
<code>a >= b</code>	Je <code>a</code> väčšie alebo rovné ako <code>b</code> ?

```

1 > 10 == 10

```

```

2 [1] TRUE
3 > 3 - 1 == 1 + 1
4 [1] TRUE
5 >
6 > 5 == 6 #je rovné?
7 [1] FALSE
8 > 5 = 6 #priradenie
9 Error in 5 = 6 : invalid (do_set) left-hand side to assignment
10 >
11 > 5 != 6
12 [1] TRUE
13 >
14 > 5 < 6
15 [1] TRUE
16 > 5 > 6
17 [1] FALSE
18 > 5 <= 6
19 [1] TRUE
20 > 5 >= 6
21 [1] FALSE

```

4.3 Logické operátory

Logické operátory podobne ako porovnávacie na výstupe dávajú logickú hodnotu TRUE alebo FALSE.

!	negácia,
&	a súčasne,
	alebo,
&&	postupné (sekvenčné) a súčasne,
	postupné (sekvenčné) alebo,
isTRUE	je logická hodnota rovná TRUE?,
%in%	operátor, pomocou ktorého identifikujeme, či prvok patrí do vektora.

```

1 > a <- 50
2 >
3 > ! (a < 100)
4 [1] FALSE
5 >
6 > a < 100
7 [1] TRUE
8 > a > 40
9 [1] TRUE
10 >

```

```
11 > a <= 30
12 [1] FALSE
13 > a >= 70
14 [1] FALSE
15 >
16 > a < 100 & a > 40
17 [1] TRUE
18 > TRUE & TRUE
19 [1] TRUE
20 >
21 > a < 100 & a <= 30
22 [1] FALSE
23 > TRUE & FALSE
24 [1] FALSE
25 >
26 > a <= 30 & a >= 70
27 [1] FALSE
28 > FALSE & FALSE
29 [1] FALSE
30 >
31 > a < 100 | a > 40
32 [1] TRUE
33 > TRUE | TRUE
34 [1] TRUE
35 >
36 > a < 100 | a <= 30
37 [1] TRUE
38 > TRUE | FALSE
39 [1] TRUE
40 >
41 > a <= 30 | a >= 70
42 [1] FALSE
43 > FALSE | FALSE
44 [1] FALSE
45 >
46 > isTRUE(a < 100)
47 [1] TRUE
48 >
49 > c("Karol", "Lukas") %in% c("Karol", "Pavol", "Jan")
50 [1] TRUE FALSE
```

4.4 Operátory, prostredníctvom ktorých vyberáme časti dátových štruktúr

O spomínaných operátoroch si viac povieme v Kapitole 6, ktorá sa venuje dátovým štruktúram. V tejto časti uvádzame ich základný prehľad.

\$	časť objektu typu data frame alebo zoznamu (listu),
[]	časť dátovej štruktúry (atomickej alebo nonatomickej),
[[]]	časť zoznamu (listu),
@	časť objektu S4.

4.5 Ďalšie používané operátory

V **R** okrem vyššie spomenutých operátorov používame množstvo ďalších, pričom uvádzame niektoré z nich. Ich použitie bude vysvetlené v nasledujúcich častiach tejto knihy.

{ }	definícia funkcií, cykly, podmienky...,
:	aritmetická postupnosť s diferenciou 1 alebo -1,
::	práca s balíkmi a funkciami,
()	poradie výpočtov, argumenty funkcie,
;	oddeľovanie výrazov pre lepšiu prehľadnosť,
#	komentár (všetko napravo v riadku je ignorované),
%>%	"pipe" operátor (pri práci s balíkmi dplyr alebo magrittr).

Kapitola 5

Funkcie

S rôznymi dátovými typmi môžeme vykonávať množstvo operácií nielen pomocou operátorov, ale môžeme využívať aj funkcie. V tejto kapitole si ukážeme použitie jednoduchých funkcií, ktoré používame pri počítaní s reálnymi číslami. Ak chceme poznať argumenty konkrétnej funkcie, môžeme na to použiť funkciu `args()`. Ukážeme si to na príklade funkcie `round()`, pomocou ktorej reálne čísla zaokrúhľujeme na vopred zadaný počet desatinných miest.

```
1 > args(round)
2 function (x, digits = 0)
3 NULL
```

Z tohto výstupu vidíme, že funkcia `round()` má dva argumenty: prvým je `x`, ktorý predstavuje číslo, ktoré chceme zaokrúhliť, a druhý argument je `digits`, ktorý predstavuje počet desatinných miest. Okrem toho vidíme aj východiskové nastavenie tejto funkcie na `digits = 0`, čo znamená, že v prípade, ak daný argument ne zadáme, funkcia zaokrúhli nami zvolené číslo na najbližšie celé číslo matematicky. Poradie argumentov, v prípade ak ich pomenujeme, nie je dôležité. Ak do funkcie priamo vpisujeme argumenty bez toho, aby sme ich pomenovali, je dôležité zachovať ich poradie. V tomto prípade prvé ide číslo, ktoré chceme zaokrúhliť a druhý je počet desatinných miest. V nasledujúcom príklade chceme číslo π zaokrúhliť na dve desatinné miesta. Ak argumenty nepomenujeme, tak je rozdiel, či použijeme funkciu `round(pi, 2)`, ktorá prvý argument (π) zaokrúhli na dve desatinné miesta alebo či použijeme funkciu `round(2, pi)`, ktorá prvý argument (2) zaokrúhli na π desatinných miest.

```
1 > pi
2 [1] 3.141593
3 >
4 > round(x = pi, digits = 2)
5 [1] 3.14
```

```

6 > round(digits = 2, x = pi)
7 [1] 3.14
8 > round(digits = 2, pi)
9 [1] 3.14
10 > round(x = pi, 2)
11 [1] 3.14
12 >
13 > round(pi, 2)
14 [1] 3.14
15 > round(2, pi)
16 [1] 2
17 >
18 > round(pi)
19 [1] 3

```

Ak uvedieme vo funkcii `round()` počet desatinných miest v argumente `digits` desatinným číslom, tak najprv sa urobí matematické zaokrúhlenie, aby sme vedeli, aký je želaný počet desatinných miest a potom sa zaokrúhli pôvodné číslo:

```

1 > pi
2 [1] 3.141593
3 > round(9.87654321, digits = pi)
4 [1] 9.877
5 > round(9.87654321, digits = 3.01)
6 [1] 9.877
7 > round(9.87654321, digits = 3.99)
8 [1] 9.8765

```

Ďalšími funkciami na zaokrúhľovanie čísel sú funkcie `floor()` na výpočet dolnej celej časti a `ceiling()` na výpočet hornej celej časti.

```

1 > floor(pi)
2 [1] 3
3 > ceiling(pi)
4 [1] 4

```

S niektorými nasledujúcimi funkciami sme sa už stretli v Kapitole 1.5. Funkcia `exp()` slúži na výpočet hodnoty Eulerovho čísla e (základu prirodzeného logaritmu) umocneného na konkrétny exponent. Ak chceme napríklad poznať hodnotu iba samotného Eulerovho čísla e použijeme argument 1 (na prvú).

```

1 > exp(1)
2 [1] 2.718282

```

Funkcia `log()` je určená na výpočet logaritmu pri ľubovoľnom základe (argument `base`), v prípade, ak tento argument nie je zadaný, je implicitne nastavené ako základ logaritmu

Eulerovo číslo e (`base = exp(1)`), vo východiskovom tvare je preto táto funkcia určená na výpočet prirodzeného logaritmu.

```
1 > log(128)
2 [1] 4.85203
3 > log(128, base = exp(1))
4 [1] 4.85203
5 > log(128, base = 2)
6 [1] 7
7 > log(1000, base = 10)
8 [1] 3
```

Funkcia `log10()` slúži na výpočet dekadického logaritmu (logaritmus so základom 10).

```
1 > log(1000, base = 10)
2 [1] 3
3 > log10(1000)
4 [1] 3
```

Už sme sa stretli aj s trigonometrickými funkciami `sin()`, `cos()` a `tan()`, kde sme upozornili na to, že je potrebné argumenty týchto funkcií uvádzať v oblúkovej a nie stupňovej miere. Vieme, že platí:

$$\pi = 180^\circ$$

$$\sin(90^\circ) = \sin\left(\frac{\pi}{2}\right) = 1$$

$$\cos(180^\circ) = \cos(\pi) = -1$$

$$\tan(45^\circ) = \tan\left(\frac{\pi}{4}\right) = 1$$

Prípadne môžeme použiť funkciu `deg2rad()` z balíka `REdaS` a priamo počítať so stupňami.

```
1 > install.packages("REdaS")
2 >
3 > library(REdaS)
4 >
5 > sin(90)
6 [1] 0.8939967
7 > sin(deg2rad(90))
8 [1] 1
9 > sin(pi / 2)
10 [1] 1
11 >
12 > cos(180)
13 [1] -0.5984601
14 > cos(deg2rad(180))
15 [1] -1
16 > cos(pi)
```

```

17 [1] -1
18 >
19 > tan(45)
20 [1] 1.619775
21 > tan(deg2rad(45))
22 [1] 1
23 > tan(pi / 4)
24 [1] 1
25 > sin(90)
26 [1] 0.8939967
27 > sin(deg2rad(90))
28 [1] 1
29 > sin(pi / 2)
30 [1] 1
31 >
32 > cos(180)
33 [1] -0.5984601
34 > cos(deg2rad(180))
35 [1] -1
36 > cos(pi)
37 [1] -1
38 >
39 > tan(45)
40 [1] 1.619775
41 > tan(deg2rad(45))
42 [1] 1
43 > tan(pi / 4)
44 [1] 1

```

Medzi jednoduché a často používané matematické funkcie môžeme zaradiť aj nasledujúce funkcie:

- `sqrt()` na výpočet druhej odmocniny konkrétneho čísla (funkcia funguje aj pre komplexné čísla, len je to potrebné pri inicializácii funkcie nadefinovať),
- `abs()` na výpočet absolútnej hodnoty konkrétneho čísla,
- `sign()` na určenie znamienka konkrétneho čísla,
- `factorial()` na výpočet faktoriálu konkrétneho čísla,
- `choose()` na výpočet kombinačného čísla.

```

1 > sqrt(25)
2 [1] 5
3 > 5 * 5
4 [1] 25

```

```
5 > 5 ^ 2
6 [1] 25
7 >
8 > sqrt(-25)
9 [1] NaN
10 Warning message:
11 In sqrt(-25) : NaNs produced
12 > sqrt(-25 + 0i)
13 [1] 0+5i
14 > (0 + 5i) ^ 2
15 [1] -25+0i
16 >
17 > abs(5)
18 [1] 5
19 > abs(-5)
20 [1] 5
21 > abs(0)
22 [1] 0
23 >
24 > sign(5)
25 [1] 1
26 > sign(-5)
27 [1] -1
28 > sign(0)
29 [1] 0
30 >
31 > factorial(5)
32 [1] 120
33 > 1 * 2 * 3 * 4 * 5
34 [1] 120
35 >
36 > choose(n = 5, k = 3)
37 [1] 10
38 > factorial(5) / (factorial(5 - 3) * factorial(3))
39 [1] 10
```

Kapitola 6

Dátové štruktúry

Dátové štruktúry v **R** možno rozdeliť do dvoch základných skupín. Prvú skupinu tvoria homogénne (atomické) dátové štruktúry, kde zaraďujeme **atomické vektory** (jednorozmerné z pohľadu dimenzionality), **atomické matice** (dvojmerné z pohľadu dimenzionality) a **atomické polia** (viacmerné z pohľadu dimenzionality). Homogénne dátové štruktúry sa vyznačujú tým, že obsahujú položky iba jedného dátového typu. Napríklad v atomickom vektore nemôžeme mať súčasne uložené numerické hodnoty a textové reťazce. Z pohľadu praktického použitia je táto homogenita značne obmedzujúca. V praxi najčastejšie analyzujeme datasety, kde v riadkoch sú jednotlivé objekty (napríklad respondenti) a v stĺpcoch sú jednotlivé premenné. Tieto premenné však môžu byť rôzneho dátového typu, napríklad pri kategorických premenných typu pohlavie alebo zamestnanie respondenta pôjde o nejaké textové reťazce (typ **character**). Naopak pri číselných spojitých premenných ako napríklad vek, výška alebo váha respondenta pôjde o numerické premenné (typ **double**, prípadne **integer**). Tieto datasety svojím usporiadaním pripomínajú atomické matice, avšak práve z dôvodu heterogenity položiek by sme ich nemohli reprezentovať touto homogénnou štruktúrou. Takéto datasety preto reprezentujeme prostredníctvom takzvaných **data frame** (dvojmerné z pohľadu dimenzionality). Jednorozmernou nehomogénnou dátovou štruktúrou je zoznam (**list**), ktorý si môžeme predstaviť tak, že napríklad ako prvý objekt v tomto zozname je nejaký atomický vektor, druhý objekt môže byť atomická matica, tretí objekt môže byť data frame, štvrtý objekt môže byť nejaká funkcia a podobne. Ide teda o akúsi množinu rôznych objektov.

6.1 Homogénne dátové štruktúry

6.1.1 Atomický vektor

Atomický vektor predstavuje základnú dátovú štruktúru v **R**. Aj keď sme doteraz vytvárali objekty, do ktorých sme priradovali iba jednu hodnotu, v skutočnosti sme vytvárali vektor dĺžky jeden, nakoľko **R** nemá definovanú dátovú štruktúru, ktorú poznáme ako skalár. Každý skalár je teda atomický vektor dĺžky jeden (obsahujúci jeden prvok). Základným princípom atomického vektora je to, že obsahuje položky iba jedného dátového typu. K vytvoreniu vektora používame funkciu `c()`: skratka z anglického *concentrate* (zlúčiť, spojiť).

Vytvoríme si objekty `vektor1`, `vektor2`, `vektor3` obsahujúce nejaké numerické hodnoty:

```
1 > vektor1 <- c(1, 2, 3, 4, 5, 6, 7, 8, 9, 10)
2 > vektor2 <- c(-1, -2, -3, -4, -5, -6, -7, -8, -9, -10)
3 > vektor3 <- c(3, 6, 9, 12, 15, 18)
4 >
5 > print(vektor1)
6 [1] 1 2 3 4 5 6 7 8 9 10
7 > print(vektor2)
8 [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10
9 > print(vektor3)
10 [1] 3 6 9 12 15 18
```

V knihe väčšinou pre zobrazenie objektov ako sú vektory, matice a podobne využívame funkciu `print()`, ktorej argumentom je názov objektu, ktorý chceme zobraziť. Tomuto zobrazeniu hovoríme *explicitné*. Okrem explicitného zobrazenia objektu poznáme aj takzvané *implicitné* zobrazenie priamo iba pomocou názvu objektu:

```
1 > print(vektor1) #explicitne zobrazenie objektu
2 [1] 1 2 3 4 5 6 7 8 9 10
3 > vektor1 #implicitne zobrazenie objektu
4 [1] 1 2 3 4 5 6 7 8 9 10
```

Ako vidíme, prvou možnosťou, ako vytvoriť vektor, je využiť funkciu `c()` a vypísať ručne všetky jeho prvky. Ak však chceme vytvoriť vektory, ktoré sme uviedli vyššie, dá sa to urobiť aj jednoduchšie. Ak chceme vytvoriť aritmetickú postupnosť s diferenciou 1, respektíve -1, to jest rad rastúcich alebo klesajúcich celých čísel, môžeme na to použiť operátor (symbol `:`). Jeho použitie je nasledovné:

```
1 > vektor1.1 <- 1:10
2 > print(vektor1.1)
3 [1] 1 2 3 4 5 6 7 8 9 10
```

Na prvý pohľad sú objekty `vektor1` a `vektor1.1` totožné. V skutočnosti je medzi nimi malý rozdiel z pohľadu dátového typu, ktorý je v nich uložený, čo si vieme overiť s využitím funkcie `typeof()`.

```
1 > typeof(vektor1)
2 [1] "double"
3 > typeof(vektor1.1)
4 [1] "integer"
```

Aritmetická postupnosť čísel (objekt `vektor1.1`) definovaná operátorom : sa automaticky považuje za rad celých čísel (typ `integer`). Vzhľadom na to, že sme objekt `vektor1` zadávali ručne, hodnoty sme definovali ako reálne čísla (typ `double`). Ak by sme ich chceli uložiť ako celé čísla (typ `integer`), museli by sme za každé z čísel písať L:

```
1 > vektor1.1.1 <- c(1L, 2L, 3L, 4L, 5L, 6L, 7L, 8L, 9L, 10L)
2 > typeof(vektor1.1.1)
3 [1] "integer"
```

Prípadne by sme mohli zmeniť triedu pôvodného vektora z typu `double` na `integer` prostredníctvom funkcie `as.integer()`:

```
1 > vektor1 <- as.integer(vektor1)
2 > typeof(vektor1)
3 [1] "integer"
```

Poznámka: V predchádzajúcom príkaze sme využili operátor priradenia nie na vytvorenie nového objektu, ale na aktualizáciu pôvodného objektu. Najprv sme hodnoty uložené v objekte `vektor1` zmenili na typ `integer` a potom sme ich priradili do objektu `vektor1`, čím sme de facto vymazali pôvodné hodnoty, ktoré boli typu `double`, a nahradili ich hodnotami typu `integer`.

Objekt `vektor2` taktiež predstavuje aritmetickú postupnosť, tentokrát klesajúcu s diferenciou rovnou -1. Na jej vytvorenie môžeme opäť použiť operátor :, prípadne môžeme využiť, že tento objekt obsahuje všetky prvky objektu `vektor1` so znamienkom mínus:

```
1 > vektor2.2 <- -1:-10
2 > print(vektor2.2)
3 [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10
4 >
5 > vektor2.2.2 <- -vektor1
6 > print(vektor2.2.2)
7 [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10
```

O tom, ako vytvárať postupnosti s inou diferenciou ako 1 alebo -1, sa dozvieme v nasledujúcej kapitole.

Postupnosti v R

Objekt `vektor3` (pozri nižšie) predstavuje aritmetickú postupnosť, jej diferenciacia však nie je -1 ani 1, preto pre jej zostavenie musíme použiť funkciu `seq()`, ktorá slúži na vytvorenie aritmetických postupností. Základnými argumentmi tejto funkcie sú:

- `from`: prvý člen aritmetickej postupnosti,
- `to`: posledný člen aritmetickej postupnosti,
- `by`: diferenciacia aritmetickej postupnosti,
- `length.out`: počet členov aritmetickej postupnosti.

Pre jednoznačnú definíciu postupnosti postačí kombinácia trojice týchto argumentov:

```
1 > seq(from = 3, to = 18, by = 3)
2 [1] 3 6 9 12 15 18
3 > seq(from = 3, to = 18, length.out = 6)
4 [1] 3 6 9 12 15 18
5 > seq(from = 3, by = 3, length.out = 6)
6 [1] 3 6 9 12 15 18
7 > seq(by = 3, to = 18, length.out = 6)
8 [1] 3 6 9 12 15 18
```

Okrem funkcie `seq()` zvykneme pri práci s vektormi používať aj iné podobné funkcie. Funkciu `seq_along` používame, ak máme nejaký vektor dĺžky `n` a chceme vytvoriť postupnosť celých čísel od 1 až po `n`:

```
1 > print(vektor3)
2 [1] 3 6 9 12 15 18
3 > length(vektor3)
4 [1] 6
5 > seq_along(vektor3)
6 [1] 1 2 3 4 5 6
```

Funkciu `seq_len()` používame, ak chceme vytvoriť postupnosť celých čísel od 1 až po `n`, argumentom tejto funkcie je teda samotné `n`, pričom ak je `n` desatinné, berie sa do úvahy jeho dolná celá časť:

```
1 > seq_len(9)
2 [1] 1 2 3 4 5 6 7 8 9
3 > seq_len(9.9)
4 [1] 1 2 3 4 5 6 7 8 9
```

Často využívanou funkciou, ak chceme vo vektore opakovávať niektoré hodnoty, je funkcia `rep()`, ktorá má tieto argumenty:

- **x**: vektor, ktorého hodnoty chceme opakovať,
- **times**: ide o prirodzené číslo alebo o vektor prirodzených čísel, ktoré definuje, koľkokrát sa hodnoty vo vektore **x** majú opakovať,
- **length.out**: prirodzené číslo, ktorým definujeme želanú dĺžku novovytvoreného vektora funkciou **rep()**,
- **each**: prirodzené číslo, ktorým vyjadrujeme, koľkokrát sa má každý argument vo vektore **x** opakovať.

```

1 > rep(c(1, 5, 8), times = 4)
2 [1] 1 5 8 1 5 8 1 5 8 1 5 8
3 >
4 > rep(c(1, 5, 8), times = c(0, 1, 2))
5 [1] 5 8 8
6 >
7 > rep(c(1, 5, 8), length.out = 8)
8 [1] 1 5 8 1 5 8 1 5
9 >
10 > rep(c(1, 5, 8), each = 3)
11 [1] 1 1 1 5 5 5 8 8 8

```

V rámci jedného použitia funkcie **rep()** môžeme kombinovať použitie týchto argumentov, napríklad:

```

1 > rep(c(1, 5, 8), each = 3, times = 4)
2 [1] 1 1 1 5 5 5 8 8 8 1 1 1 5 5 5 8 8 8 1 1 1 5 5 5 8 8 8 1 1 1 5
3 [32] 5 5 8 8 8
4 >
5 > rep(c(1, 5, 8), times = 4, length.out = 10)
6 [1] 1 5 8 1 5 8 1 5 8 1
7 >
8 > rep(c(1, 5, 8), each = 4, length.out = 10)
9 [1] 1 1 1 1 5 5 5 5 8 8
10 >
11 > rep(c(1, 5, 8), times = 2)
12 [1] 1 5 8 1 5 8
13 >
14 > rep(c(1, 5, 8), times = 2, length.out = 10)
15 [1] 1 5 8 1 5 8 1 5 8 1

```

Všimnime si, že argument **length.out = 10** spôsobí, že aj keď by tri hodnoty pri dvojnásobnom opakovaní mali vytvoriť vektor dĺžky šesť, novovzniknutý vektor bude mať dĺžku desať, pričom hodnoty na 7. až 10. mieste sú opakujúce sa hodnoty zo začiatku vektora (7. hodnota = 1. hodnota, 8. hodnota = 2. hodnota a podobne).

Matematické operácie s vektormi

Ak používame matematické operátory vo vzťahu skalár a vektor, táto matematická operácia sa vykoná osobitne pre skalár a každý prvok vektora, napríklad násobenie skalárom (možnosť ako alternatívne definovať objekt `vektor3`):

```
1 > 3 * 1:6
2 [1]  3  6  9 12 15 18
```

Poznámka: V predchádzajúcom príkaze sme najprv vytvorili vektor s aritmetickou postupnosťou čísel od 1 do 6 a následne sme každý prvok tohto vektora vynásobili číslom 3.

Príklad iných matematických operácií vo vzťahu skalár a vektor uvádzame v nasledujúcej časti:

```
1 > print(vektor3)
2 [1]  3  6  9 12 15 18
3 >
4 > vektor3 / 2
5 [1] 1.5 3.0 4.5 6.0 7.5 9.0
6 >
7 > vektor3 ^ 2
8 [1]  9  36  81 144 225 324
9 >
10 > vektor3 %% 2 #zvysok po deleni 2
11 [1] 1 0 1 0 1 0
12 >
13 > vektor3 %/% 2 #celociselne delenie 2
14 [1] 1 3 4 6 7 9
```

Predtým uvádzané funkcie z Kapitoly 5 vo vzťahu k vektorom fungujú na rovnakom princípe, to znamená, že konkrétna funkcia sa aplikuje na každý jeden prvok vektora, pričom funkcie môžeme aj vnárať do seba, napríklad:

```
1 > print(vektor1)
2 [1]  1  2  3  4  5  6  7  8  9 10
3 >
4 > sqrt(vektor1)
5 [1] 1.000000 1.414214 1.732051 2.000000 2.236068 2.449490
6 [7] 2.645751 2.828427 3.000000 3.162278
7 >
8 > round(sqrt(vektor1), digits = 2)
9 [1] 1.00 1.41 1.73 2.00 2.24 2.45 2.65 2.83 3.00 3.16
```

Pri práci s funkciami môžeme vektory využívať aj ako argumenty týchto funkcií. Ukážeme si to na príklade funkcie `choose()`, ktorú používame pri počítaní kombinačných čísel. Chceli by sme napríklad poznať hodnoty:

$$\binom{10}{1}, \binom{10}{2}, \dots, \binom{10}{10}$$

Prvý argument `n` funkcie `choose()` sa nemení a stále je rovný 10, druhý argument `k` nadobúda hodnoty uložené v objekte `vektor1`:

```
1 > choose(n = 10, k = vektor1)
2 [1] 10 45 120 210 252 210 120 45 10 1
```

Funkcie vhodné na prácu s vektormi

Existuje množstvo ďalších funkcií, ktoré sú určené na prácu s celými vektormi a nevyhodnocujú sa po jeho jednotlivých zložkách. Napríklad funkcia `sum()` vypočíta súčet všetkých hodnôt vektora a na výstupe vypíše iba jednu hodnotu:

```
1 > print(vektor3)
2 [1] 3 6 9 12 15 18
3 > sum(vektor3)
4 [1] 63
```

Dĺžkou vektora sa rozumie počet prvkov, ktoré vektor obsahuje, preto ak chceme poznať tento počet, využívame na to funkciu `length()`:

```
1 > length(vektor1)
2 [1] 10
3 > length(vektor2)
4 [1] 10
5 > length(vektor3)
6 [1] 6
```

Základné popisné charakteristiky hodnôt numerického vektora (minimum, maximum, kvartily, aritmetický priemer, medián) vieme vypísať prostredníctvom funkcie `summary()`:

```
1 > print(vektor1)
2 [1] 1 2 3 4 5 6 7 8 9 10
3 > summary(vektor1)
4   Min. 1st Qu.  Median    Mean 3rd Qu.    Max.
5   1.00   3.25   5.50   5.50   7.75   10.00
6 > min(vektor1)
7 [1] 1
8 > quantile(vektor1, 1 / 4)
9 25%
10 3.25
11 > median(vektor1) #alebo quantile(vektor1, 1 / 2)
12 [1] 5.5
13 > mean(vektor1)
14 [1] 5.5
```

```

15 > quantile(vektor1, 3 / 4)
16 75%
17 7.75
18 > max(vektor1)
19 [1] 10

```

Medzi ďalšie používané funkcie patria:

- `range()`: dolná hranica (minimum) a horná hranica (maximum),
- `IQR()`: medzikvartilové rozpätie,
- `prod()`: súčin hodnôt vektora,
- `sort()`: vzostupné (`decreasing = TRUE`) alebo zostupné (`decreasing = FALSE`) usporiadanie hodnôt vektora,
- `cumsum()`: kumulatívny súčet hodnôt vektora,
- `sd()`: výberová smerodajná odchýlka hodnôt vektora (druhá odmocnina výberového rozptylu),
- `var()`: výberový rozptyl hodnôt vektora (druhá mocnina výberovej smerodajnej odchýlky).

```

1 > print(vektor3)
2 [1] 3 6 9 12 15 18
3 > IQR(vektor3)
4 [1] 7.5
5 > range(vektor3)
6 [1] 3 18
7 > prod(vektor3)
8 [1] 524880
9 > sort(c(1, 8, 6, 4), decreasing = TRUE)
10 [1] 8 6 4 1
11 > sort(c(1, 8, 6, 4), decreasing = FALSE)
12 [1] 1 4 6 8
13 > cumsum(vektor3)
14 [1] 3 9 18 30 45 63
15 > sd(vektor3)
16 [1] 5.612486
17 > sqrt(var(vektor3))
18 [1] 5.612486
19 > var(vektor3)
20 [1] 31.5
21 > sd(vektor3) ^ 2
22 [1] 31.5

```

Pomocou funkcie `c()` je možné nielen vytvoriť nový vektor, ale taktiež spojiť viacero vektorov dokopy, prípadne pridať nejaké hodnoty na začiatok alebo na koniec vektora:

```
1 > print(vektor3)
2 [1] 3 6 9 12 15 18
3 > novyVektor3 <- c(0, vektor3, 21, 24)
4 > print(novyVektor3)
5 [1] 0 3 6 9 12 15 18 21 24
```

```
1 > print(vektor2)
2 [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10
3 > print(vektor3)
4 [1] 3 6 9 12 15 18
5 > c(vektor2, vektor3)
6 [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10 3 6 9 12 15 18
```

Princíp automatickej konverzie hodnôt vektora na nadradený dátový typ

V úvode kapitoly sme si uviedli, že atomický vektor je homogénna dátová štruktúra, ktorej všetky hodnoty sú rovnakého typu (napríklad celé číslo: typ `integer` alebo textový reťazec typ `character` a podobne). Ak by sme predsa len do jedného vektora chceli uložiť rôzne dátové typy, je možné to urobiť, avšak funkcia `c()` vykoná automatickú konverziu na nadradený dátový typ alebo na taký typ, na ktorý je možné konvertovať všetky hodnoty bez straty informácie. V praxi to napríklad znamená, že ak sme mali vektor logických hodnôt typu `logical` a pridali by sme k nemu nejaké celé číslo (typ `integer`), automatická konverzia premení všetky hodnoty `TRUE` na `1L` a všetky hodnoty `FALSE` na `0L`:

```
1 > vektorLogickychHodnot <- c(TRUE, FALSE, TRUE, TRUE)
2 > print(vektorLogickychHodnot)
3 [1] TRUE FALSE TRUE TRUE
4 > typeof(vektorLogickychHodnot)
5 [1] "logical"
6 >
7 > konverze1 <- c(vektorLogickychHodnot, 5L)
8 > print(konverze1)
9 [1] 1 0 1 1 5
10 > typeof(konverze1)
11 [1] "integer"
```

Ak by sme v pôvodnom vektore mali iba celé čísla a pridáme tam nejaké desatinné číslo, automaticky všetky hodnoty vo vektore už nebudú typu `integer`, ale konvertujú sa na typ `double`:

```
1 > print(konverze1)
```

```

2 [1] 1 0 1 1 5
3 > typeof(konverze1)
4 [1] "integer"
5 > konverze2 <- c(konverze1, 3.8)
6 > print(konverze2)
7 [1] 1.0 0.0 1.0 1.0 5.0 3.8
8 > typeof(konverze2)
9 [1] "double"

```

Ak by sme v pôvodnom vektore mali iba desatinné čísla a pridáme tam nejaký textový reťazec, automaticky všetky hodnoty vo vektore už nebudú typu `double`, ale konvertujú sa na typ `character`:

```

1 > print(konverze2)
2 [1] 1.0 0.0 1.0 1.0 5.0 3.8
3 > typeof(konverze2)
4 [1] "double"
5 > konverze3 <- c(konverze2, "ahoj")
6 > print(konverze3)
7 [1] "1"      "0"      "1"      "1"      "5"      "3.8"    "ahoj"
8 > typeof(konverze3)
9 [1] "character"

```

Treba upozorniť, že touto konverziou sa pôvodné číselné hodnoty dajú do úvodzoviek a pracuje sa s nimi ďalej ako s textovým reťazcom, teda nie je možné už s nimi ďalej vykonávať numerické operácie. Konverzia základných dátových typov je teda nasledovná:

logical → integer → double → character

Pomenovanie prvkov vektora

Prvky vektora môžeme aj pomenovať, na čo najčastejšie využívame funkciu `names()`. Okrem toho môžeme prvky vektora pomenovať už priamo pri inicializácii funkcie `c()` alebo pomocou funkcie `attr()`. Ukážeme si všetky tri spôsoby, pričom názvy jednotlivých prvkov budú písmená latinskej abecedy, pri ich generovaní využijeme funkciu `LETTERS()`: funkcia na generovanie veľkých písmen latinskej abecedy, existuje aj funkcia `letters()`: funkcia na generovanie malých písmen latinskej abecedy:

```

1 > LETTERS
2 [1] "A" "B" "C" "D" "E" "F" "G" "H" "I" "J" "K" "L" "M" "N" "O"
3 [16] "P" "Q" "R" "S" "T" "U" "V" "W" "X" "Y" "Z"
4 > letters
5 [1] "a" "b" "c" "d" "e" "f" "g" "h" "i" "j" "k" "l" "m" "n" "o"
6 [16] "p" "q" "r" "s" "t" "u" "v" "w" "x" "y" "z"
7 > LETTERS[3]
8 [1] "C"

```

```

9 > LETTERS[1:5]
10 [1] "A" "B" "C" "D" "E"

```

```

1 > print(vektor1)
2 [1] 1 2 3 4 5 6 7 8 9 10
3 > length(vektor1)
4 [1] 10
5 > names(vektor1) <- LETTERS[1:10]
6 > print(vektor1)
7 A B C D E F G H I J
8 1 2 3 4 5 6 7 8 9 10

```

Ak chceme vypísať hodnoty vektora, ktoré sú pomenované, ale tieto názvy nechceme vo výpise vidieť, môžeme na to použiť funkciu `unname()`:

```

1 > print(vektor1)
2 A B C D E F G H I J
3 1 2 3 4 5 6 7 8 9 10
4 > print(unname(vektor1))
5 [1] 1 2 3 4 5 6 7 8 9 10

```

Následne je možné sa týmto názvom prvku aj odkazovať na konkrétny prvok vektora:

```

1 > vektor1["D"]
2 D
3 4

```

```

1 > print(vektor3)
2 [1] 3 6 9 12 15 18
3 > length(vektor3)
4 [1] 6
5 > attr(vektor3, "names") <- LETTERS[1:6]
6 > print(vektor3)
7 A B C D E F
8 3 6 9 12 15 18

```

```

1 > vektor4 <- c("jmeno1" = 3, "jmeno2" = 4, "jmeno3" = 5)
2 > print(vektor4)
3 jmeno1 jmeno2 jmeno3
4      3      4      5

```

Pri vytvorení objektu `vektor4` (viď vyššie) sme hneď pri jeho vytváraní definovali pomenovanie jeho jednotlivých hodnôt. Tieto názvy sme dávali do úvodzoviek.

Vyberanie niektorých hodnôt vektora

Ak chceme pracovať iba s niektorými hodnotami vektora, využívame na to takzvané indexovanie, na čo používame hranaté zátvorky. Programovací jazyk **R**, na rozdiel napríklad od programovacieho jazyka Python, nepoužíva tzv. *zero base indexing*, pri ktorom hodnoty na prvej pozícii majú index (poradie) nula, ale hodnotu jeden. V prostredí **R**, ak chceme odkazovať na k -tú hodnotu, použijeme na to hranatú zátvorku s indexom k .

Chceme pracovať s treťou hodnotou objektu `vektor1`:

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > vektor1[3]
5 C
6 3
```

Ak chceme pracovať s hodnotami nasledujúcimi bezprostredne za sebou, môžeme si pomôcť operátorom aritmetickej postupnosti s diferenciou 1 alebo -1 (:), napríklad ak chceme druhú až piatu hodnotu objektu `vektor1`:

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > vektor1[2:5]
5 B C D E
6 2 3 4 5
```

Ak chceme pracovať s hodnotami, ktoré nenasledujú bezprostredne za sebou, nemôžeme to urobiť len vypísaním poradí do hranatých zátvoriek, nakoľko v takomto prípade vyberáme prvky na základe dimenzií a vektor je iba jednodimenzionálny objekt. Ak chceme pracovať s prvou a piatou hodnotou objektu `vektor1`, nasledujúca syntax nie je správna:

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > vektor1[1, 5]
5 Error in vektor1[1, 5] : incorrect number of dimensions
```

Vidíme, že chybová hláška nás upozorňuje, že počet dimenzií nie je korektný. Ak by sme pracovali s nejakým dvojrozmerným objektom (matica alebo data frame), pomocou tejto syntaxe by sme dostali hodnotu nachádzajúcu sa v prvom riadku a piatom stĺpci. Ak však chceme pracovať s prvou a piatou hodnotou objektu `vektor1`, musíme pre zadefinovanie týchto poradí použiť funkciu `c()`:

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
```

```

3  1  2  3  4  5  6  7  8  9 10
4  > vektor1[c(1, 5)]
5  A E
6  1 5

```

Odkazovať na prvky vektora je možné aj pomocou vektora logických hodnôt. Ak máme vektor číselných hodnôt dĺžky n a nejaký vektor logických hodnôt rovnakej dĺžky, potom ak za vektor číselných hodnôt dáme hranaté zátvorky, do ktorých vpíšeme názov vektora logických hodnôt, dostaneme hodnoty vektora číselných hodnôt na tých poradiach, pre ktoré sú hodnoty vo vektore logických hodnôt rovné TRUE:

```

1  > vektor5 <- seq(from = 2, by = 7, length.out = 4)
2  > print(vektor5)
3  [1]  2  9 16 23
4  > print(vektorLogickychHodnot)
5  [1] TRUE FALSE TRUE TRUE
6  > vektor5[vektorLogickychHodnot]
7  [1]  2 16 23

```

Na tomto princípe potom môžeme vykonávať rôzne filtrovanie údajov. Jednoducho vytvoríme nejaký vektor logických hodnôt, ktorý napríklad odpovedá na otázku, či daná hodnota vo vektore je kladná, deliteľná nejakým číslom, väčšia ako priemer a podobne. Následne pomocou hranatých zátvoriek budeme pracovať iba s hodnotami, ktoré danú vlastnosť spĺňajú. Praktický príklad takéhoto jednoduchého filtrovania uvádzame v Kapitole 8.2.

Uvažujme, že chceme pracovať iba s párnymi (v ČJ sudými) číslami z objektu `vektor3`. Na overenie deliteľnosti využívame aritmetický operátor zvyšku po delení (modulo: `%%`):

```

1  > print(vektor3)
2  A  B  C  D  E  F
3  3  6  9 12 15 18
4  > vektor3 %% 2 == 0
5      A      B      C      D      E      F
6 FALSE TRUE FALSE TRUE FALSE TRUE
7 > vektor3sude <- vektor3 %% 2 == 0
8 > print(vektor3sude)
9      A      B      C      D      E      F
10 FALSE TRUE FALSE TRUE FALSE TRUE
11 > vektor3[vektor3sude]
12  B  D  F
13  6 12 18

```

Filtrovanie môžeme robiť aj bez toho, aby sme vytvárali vektor logických hodnôt. Predstavme si, že chceme napríklad v objekte `vektor1` pracovať iba s hodnotami, ktoré sú väčšie ako priemer:

```

1 > print(vektor1)
2   A  B  C  D  E  F  G  H  I  J
3   1  2  3  4  5  6  7  8  9 10
4 > mean(vektor1)
5 [1] 5.5
6 > vektor1[vektor1 > mean(vektor1)]
7   F  G  H  I  J
8   6  7  8  9 10

```

Tento filter môžeme skombinovať aj s inými logickými operátormi, napríklad s operátorom logického súčtu (a súčasne: &). Uvažujme, že v objekte vektor1 chceme pracovať iba s hodnotami, ktoré sú väčšie ako priemer a súčasne sú deliteľné tromi:

```

1 > print(vektor1)
2   A  B  C  D  E  F  G  H  I  J
3   1  2  3  4  5  6  7  8  9 10
4 > vektor1[vektor1 > mean(vektor1) & vektor1 %% 3 == 0]
5 F I
6 6 9

```

Pomocou hranatých zátvoriek môžeme nielen vyberať niektoré hodnoty vektora, ale tak tiež môžeme povedať, ktoré hodnoty nechceme, aby boli vybraté. Tieto hodnoty vo vektore „nevyberieme“ prostredníctvom operátora - :

```

1 > print(vektor5)
2 [1]  2  9 16 23
3 > vektor5[-3]
4 [1]  2  9 23
5 > vektor5[-c(1, 4)]
6 [1]  9 16

```

Práca s viacerými vektormi

Už sme si uviedli, že aritmetické a logické operácie v kombinácii skalár (vektor dĺžky jeden) a vektor sa realizujú tak, že pre každý prvok vektora vykonáme túto operáciu so skalárom osobitne:

```

1 > 1:6
2 [1] 1 2 3 4 5 6
3 > 3 * 1:6
4 [1]  3  6  9 12 15 18

```

```

1 > c(TRUE, FALSE, TRUE)
2 [1] TRUE FALSE TRUE

```

```

3 > TRUE | c(TRUE, FALSE, TRUE)
4 [1] TRUE TRUE TRUE

```

Ak pracujeme s viacerými vektormi, všetky aritmetické a logické operácie sa vo vektoroch prevádzajú po prvkoch:

```

1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > length(vektor1)
5 [1] 10
6 > print(vektor2)
7  [1] -1 -2 -3 -4 -5 -6 -7 -8 -9 -10
8 > length(vektor2)
9 [1] 10
10 >
11 > vektor1 + vektor2
12 A B C D E F G H I J
13 0 0 0 0 0 0 0 0 0 0
14 > vektor1 * vektor2
15  A  B  C  D  E  F  G  H  I  J
16 -1 -4 -9 -16 -25 -36 -49 -64 -81 -100

```

Ak sa dĺžka vektorov odlišuje, **R** automaticky opakuje jeho hodnoty znovu a znovu a upozorní na to po realizácii danej operácie.

```

1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > length(vektor1)
5 [1] 10
6 > print(vektor3)
7  A  B  C  D  E  F
8  3  6  9 12 15 18
9 > length(vektor3)
10 [1] 6
11 >
12 > vektor1 + vektor3
13  A  B  C  D  E  F  G  H  I  J
14  4  8 12 16 20 24 10 14 18 22
15 Warning message:
16 In vektor1 + vektor3 :
17  longer object length is not a multiple of shorter object length

```

```

1 > print(vektorLogickyHodnot)
2 [1] TRUE FALSE TRUE TRUE

```

```

3 > length(vektorLogickychHodnot)
4 [1] 4
5 > print(vektor3sude)
6      A      B      C      D      E      F
7 FALSE TRUE FALSE TRUE FALSE TRUE
8 > length(vektor3sude)
9 [1] 6
10 > vektorLogickychHodnot & vektor3sude
11      A      B      C      D      E      F
12 FALSE FALSE FALSE TRUE FALSE FALSE
13 Warning message:
14 In vektorLogickychHodnot & vektor3sude :
15 longer object length is not a multiple of shorter object length

```

R však varovanie nevydá v prípade, ak dĺžka kratšieho vektora je celočíselným deliteľom dĺžky dlhšieho vektora:

```

1 > print(vektor1)
2  A B  C  D E F G H  I  J
3  1 2  3  4 5 6 7 8  9 10
4 > length(vektor1)
5 [1] 10
6 > print(konverze1)
7 [1] 1 0 1 1 5
8 > length(konverze1)
9 [1] 5
10 >
11 > vektor1 + konverze1
12  A B  C  D E F G H  I  J
13  2 2  4  5 10 7 7 9 10 15

```

Skalárny súčin vektorov

V praxi sa často stretávame so situáciou, keď potrebujeme vypočítať takzvaný *skalárny súčin dvoch vektorov*. Nech $\vec{u} = (u_1, u_2, \dots, u_n)$ a $\vec{v} = (v_1, v_2, \dots, v_n)$ sú vektory. Skalárny súčin vektorov $\vec{u}$ a $\vec{v}$ je definovaný nasledovne:

$$\vec{u} \cdot \vec{v} = u_1 v_1 + u_2 v_2 + \dots + u_n v_n \quad (6.1)$$

Ako sme si ukázali, operátor súčinu ($*$) vo vektoroch funguje tak, že vynásobí navzájom odpovedajúce si hodnoty a výsledkom je zase vektor rovnakého rozmeru. Skalárny súčin však predstavuje súčet týchto súčinov, takže ak by sme ho chceli vypočítavať, mohli by sme postupovať nasledovne:

```

1 > print(vektor1)
2   A  B  C  D  E  F  G  H  I  J
3   1  2  3  4  5  6  7  8  9 10
4 > print(vektor2)
5   [1]  -1  -2  -3  -4  -5  -6  -7  -8  -9 -10
6 > vektor1 * vektor2
7     A    B    C    D    E    F    G    H    I    J
8    -1   -4   -9  -16  -25  -36  -49  -64  -81 -100
9 > sum(vektor1 * vektor2)
10 [1] -385

```

Existuje aj jednoduchší spôsob, pri ktorom využívame operátor, ktorý sa najčastejšie používa, keď chceme vypočítať súčin dvoch matíc. Tento operátor označujeme `%*%` a jeho použitie je nasledovné:

```

1 > vektor1 %*% vektor2
2     [,1]
3 [1,] -385

```

Pri skalárnom súčine je ale nevyhnutné, aby oba vektory boli rovnakej dĺžky, inak operátor maticového súčinu nie je možné použiť a na výstupe dostaneme chybovú hlášku, s ktorou sa stretávame aj v prípade, ak sa snažíme násobiť matice, ktoré nemajú správny rozmer.

```

1 > print(vektor1)
2   A  B  C  D  E  F  G  H  I  J
3   1  2  3  4  5  6  7  8  9 10
4 > length(vektor1)
5 [1] 10
6 > print(vektor3)
7   A  B  C  D  E  F
8   3  6  9 12 15 18
9 > length(vektor3)
10 [1] 6
11 > vektor1 %*% vektor3
12 Error in vektor1 %*% vektor3 : non-conformable arguments

```

Vektory generujúce náhodné čísla

Niekedy potrebujeme vygenerovať vektory s náhodnými číslami, aby sme na nich mohli aplikovať nejaké výpočty a podobne. Existuje viacero funkcií na generovanie pseudonáhodných čísel v **R**. Prvou funkciou je funkcia `runif()`, ktorú používatelia Excelu poznajú ako **NÁHČÍSLO**, respektíve **RAND**. Najdôležitejším argumentom funkcie je `n`, ktorý predstavuje počet pseudonáhodných čísel, ktoré chceme vygenerovať. Pokiaľ nezadefinujeme žiadne ďalšie argumenty, funkcia `runif()` generuje `n` rovnomerne rozdelených pseudoná-

hodných čísel z intervalu $\langle 0, 1 \rangle$. Ak chceme generovať n pseudonáhodných čísel z intervalu $\langle a, b \rangle$, môžeme pridať argumenty funkcie `min` a `max`:

```
1 > runif(n = 5)
2 [1] 0.46144747 0.15706526 0.05461457 0.16340385 0.01052053
3 > runif(5)
4 [1] 0.2384374 0.4018661 0.4601084 0.7333281 0.6895014
5 > runif(n = 5, min = 1, max = 3)
6 [1] 1.986712 2.218669 2.384245 2.105338 1.659767
7 > runif(n = 5, min = 1, max = 3)
8 [1] 1.079947 1.169818 1.435591 1.403885 2.269265
```

V predchádzajúcom kóde sme dvakrát za sebou vygenerovali päť pseudonáhodných čísel, prvýkrát z intervalu $\langle 0, 1 \rangle$ a druhýkrát z intervalu $\langle 1, 3 \rangle$. Ako môžeme vidieť, hodnoty sú naozaj generované náhodne a pri každom spustení kódu sú rôzne. V praxi sa však stretávame aj so situáciou, keď potrebujeme generovať rovnaké náhodné čísla, aby sme si napríklad mohli overiť výsledky nejakých simulácií a podobne. Na to slúži funkcia `set.seed()`, ktorá predstavuje doslova akési „semienko“, ktoré sa využíva pri inicializácii generátora pseudonáhodných čísel. Tento generátor pri použití konkrétneho „semienka“ stále generuje rovnakú sekvenciu pseudonáhodných čísel. Nastavme si toto „semienko“ napríklad na hodnotu 123 a skúsme vygenerovať tieto pseudonáhodné čísla opäť:

```
1 > set.seed(123)
2 > runif(n = 5, min = 1, max = 3)
3 [1] 1.575155 2.576610 1.817954 2.766035 2.880935
4 > set.seed(123)
5 > runif(n = 5, min = 1, max = 3)
6 [1] 1.575155 2.576610 1.817954 2.766035 2.880935
```

S takýmto nastavením „semienka“ dostaneme na ľubovoľnom počítači pri generovaní päť pseudonáhodných čísel z intervalu $\langle 1, 3 \rangle$ práve týchto päť hodnôt.

Ďalšou funkciou, ktorá nám umožňuje generovať napríklad náhodné celé čísla, je funkcia `sample()`. Podstatou tejto funkcie je, že na vstupe zadáme vektor s ľubovoľnými hodnotami (argument `x`) a povieme koľko hodnôt chceme z tohto vektora náhodne vytiahnuť (argument `size`), pričom môžeme uvažovať ako výbery s opakovaním, tak aj výbery bez opakovania (argument `replace`). Ak by sme chceli napríklad vybrať náhodne päť celých čísel z intervalu 1 až 10, pričom tieto čísla sa môžu opakovať, postupovali by sme nasledovne:

```
1 > sample(x = 1:10,
2 +       size = 5,
3 +       replace = TRUE)
4 [1] 7 9 9 10 7
```

Vidíme, že ak máme nastavený argument `replace` na logickú hodnotu `TRUE`, ide o tzv.

výbery s opakovaním, čo znamená, že hodnoty sa môžu opakovať (v našom prípade sa opakujú sedmičky a deviatky). Ak nechceme výbery s opakovaním, nastavíme tento argument na `FALSE`:

```
1 > sample(x = 1:10,
2 +       size = 5,
3 +       replace = FALSE)
4 [1]  5  7 10  3  6
```

V tejto časti je potrebné upozorniť na časté chyby spôsobené tým, že z vektora dĺžky `n` chceme náhodne vybrať viac ako `n` hodnôt. V takomto prípade je daný výber logicky možný iba tak, že hodnoty sa budú opakovať. Preto funkcia `sample()` s argumentom `replace = FALSE` v takomto prípade vypisuje chybovú hlášku:

```
1 > sample(x = 1:10,
2 +       size = 12,
3 +       replace = FALSE)
4 Error in sample.int(length(x), size, replace, prob) :
5   cannot take a sample larger than the population when 'replace =
   FALSE'
```

Funkcia `sample()` má veľmi univerzálne použitie. Argument `x` môže byť vektor s údajmi ľubovoľného dátového typu, napríklad aj `character`: textový reťazec.

V objekte `state.name` nachádzajúcim sa priamo v balíčku `datasets`, ktorý sa automaticky načíta pri spustení **R**, máme uvedený abecedný zoznam všetkých štátov USA. Ide o vektor dĺžky 50 typu `character`:

```
1 > typeof(state.name)
2 [1] "character"
```

Ak z tohto vektora chceme náhodne vybrať 10 štátov bez opakovania, vykonáme to nasledovne:

```
1 > sample(x = state.name,
2 +       size = 10,
3 +       replace = FALSE)
4 [1] "Michigan"      "Wyoming"      "Rhode Island" "New Mexico"
5 [5] "Kentucky"     "North Dakota" "Arkansas"     "Illinois"
6 [9] "California"   "Missouri"
7 > sample(x = state.name,
8 +       size = 10,
9 +       replace = FALSE)
10 [1] "New Mexico"    "Virginia"     "New Jersey"   "Ohio"
11 [5] "Indiana"       "New Hampshire" "New York"     "Connecticut"
12 [9] "Arizona"      "Minnesota"
```


Ako môžeme vidieť, výber je stále náhodný, opäť však máme možnosť použiť funkciu `set.seed()` a nastaviť generátor pseudonáhodných čísel tak, aby nám stále generoval rovnaký zoznam štátov:

```
1 > set.seed(123)
2 > sample(x = state.name,
3 +       size = 10,
4 +       replace = FALSE)
5 [1] "New Mexico"      "Iowa"             "Indiana"          "Arizona"
6 [5] "Tennessee"       "Texas"            "Oregon"           "West Virginia"
7 [9] "Missouri"        "Montana"
8 >
9 > set.seed(123)
10 > sample(x = state.name,
11 +       size = 10,
12 +       replace = FALSE)
13 [1] "New Mexico"      "Iowa"             "Indiana"          "Arizona"
14 [5] "Tennessee"       "Texas"            "Oregon"           "West Virginia"
15 [9] "Missouri"        "Montana"
```

Funkcia `sample()` má ešte jeden argument s názvom `prob`, prostredníctvom ktorého môžeme nadefinovať pravdepodobnosti, s ktorými sa jednotlivé hodnoty vektora definovaného argumentom `x` budú vyberať. Povedzme, že chceme vytvoriť objekt s názvom `hodnoty`, ktorý bude obsahovať 100 náhodne vygenerovaných celých čísel z intervalu $\langle 1, 4 \rangle : \{1, 2, 3, 4\}$, pričom pravdepodobnosť výberu jednotky je 0,1, pravdepodobnosť výberu dvojky je 0,2, pravdepodobnosť výberu trojky je 0,3 a pravdepodobnosť výberu štvorky je 0,4. Ak sa niektoré hodnoty vo vektore opakujú, je vhodné ho reprezentovať jednoduchou tabuľkou početností, využívame pri tom funkciu `table()`:

```
1 > hodnoty <- sample(
2 +   x = 1:4,
3 +   size = 100,
4 +   replace = TRUE,
5 +   prob = c(0.1, 0.2, 0.3, 0.4)
6 + )
7 >
8 > print(hodnoty)
9 [1] 2 4 3 2 1 4 1 2 3 4 4 3 3 3 1 3 3 3 4 4 4 4 2 4 2 3 3 4 3 4
10 [31] 2 4 1 1 2 4 4 3 4 3 2 4 3 3 2 1 2 3 1 3 3 4 4 4 3 2 4 4 4 2
11 [61] 2 1 3 4 3 2 4 4 4 4 4 4 4 4 3 4 4 4 2 2 2 1 4 4 2 2 4 2 2 3
12 [91] 3 4 4 3 3 4 3 2 4 4
13 >
14 > table(hodnoty)
15 hodnoty
16 1  2  3  4
17 9 22 27 42
```

Ako môžeme vidieť, argument `prob` ovplyvnil, že rozdelenie hodnôt nie je rovnomerné a má veľmi blízko vektoru pravdepodobností, ktorý sme nadefinovali. Ak by sme daný argument nepoužili, rozdelenie hodnôt by bolo rovnomerné, nakoľko každá z hodnôt 1 až 4 by mala rovnakú pravdepodobnosť, že bude vybraná.

```
1 > hodnoty2 <- sample(x = 1:4,
2 +                   size = 100,
3 +                   replace = TRUE)
4 >
5 > print(hodnoty2)
6 [1] 1 4 1 1 3 3 1 1 3 1 3 4 2 1 1 3 4 2 3 2 1 2 3 2 3 3 1 2 4 1
7 [31] 2 2 4 2 4 2 1 4 2 4 1 4 2 1 4 4 3 2 1 3 2 4 3 2 3 3 1 4 2 3
8 [61] 1 1 3 2 3 2 2 2 4 1 4 2 1 4 4 1 4 1 4 2 2 4 3 3 1 2 4 4 2 3
9 [91] 3 3 3 2 1 4 1 4 3 4
10 >
11 > table(hodnoty2)
12 hodnoty2
13 1 2 3 4
14 25 26 24 25
```

Existujú aj iné funkcie, ktoré generujú náhodné hodnoty napríklad z vybraných rozdelení pravdepodobností. Tieto funkcie začínajú písmenom `r`, za ktorým nasleduje skratka pre konkrétne rozdelenie, napríklad ak chceme generovať náhodné hodnoty z normálneho rozdelenia, využijeme funkciu `rnorm()`, ak chceme generovať náhodné hodnoty z binomického rozdelenia, využijeme funkciu `rbinom()` a podobne. Argumenty funkcií závisia od parametrov konkrétnych rozdelení a v tomto učebnom materiáli sa im nebudeme viac venovať.

Ďalšie užitočné funkcie

Najmä pri práci s maticami a datasetmi (neatomická štruktúra `data frame`) je veľmi vhodné vypísať si iba prvých, respektíve posledných pár pozorovaní (riadkov), aby sme zistili, s akými dátami pracujeme a podobne. Na to používame funkcie `head()` a `tail()`, ktoré je možné aplikovať aj na vektory. Namiesto toho, aby sme vypisovali celý vektor so sto hodnotami, môžeme si nechať vypísať iba prvých alebo posledných pár hodnôt vektora (prednastavené na prvých, respektíve posledných šesť hodnôt: možno zmeniť argumentom `n`).

```
1 > print(hodnoty)
2 [1] 2 4 3 2 1 4 1 2 3 4 4 3 3 3 1 3 3 3 4 4 4 4 2 4 2 3 3 4 3 4
3 [31] 2 4 1 1 2 4 4 3 4 3 2 4 3 3 2 1 2 3 1 3 3 4 4 4 3 2 4 4 4 2
4 [61] 2 1 3 4 3 2 4 4 4 4 4 4 4 4 3 4 4 4 2 2 2 1 4 4 2 2 4 2 2 3
5 [91] 3 4 4 3 3 4 3 2 4 4
6 > head(hodnoty)
7 [1] 2 4 3 2 1 4
```

```

8 > tail(hodnoty)
9 [1] 3 4 3 2 4 4
10 > tail(hodnoty, n = 8)
11 [1] 4 3 3 4 3 2 4 4

```

Ak chceme otestovať, či objekt je atomický vektor, používame funkciu `is.atomic()`, ktorá vráti logickú hodnotu `TRUE` v prípade, že všetky prvky vektora sú rovnakého dátového typu (sú homogénne), a `FALSE` v prípade, že ide o neatomickú štruktúru `list` (zoznam) alebo `data frame`. Neatomický vektor: `list` (zoznam), môže obsahovať prvky rôzneho dátového typu (nehomogénne). Existuje aj funkcia `is.vector()`, ktorá však aj v prípade atomického, aj neatomického vektora vracia logickú hodnotu `TRUE`.

```

1 > print(konverze3)
2 [1] "1"      "0"      "1"      "1"      "5"      "3.8"    "ahoj"
3 > is.atomic(konverze3)
4 [1] TRUE

```

Inicializácia vektora

V praktických príkladoch najmä pri programovaní cyklov často na začiatku inicializujeme (vytvárame) vektor daného typu a konkrétnej dĺžky. Robíme to z dôvodu eliminácie možných chýb, ktoré by mohli vzniknúť priradením nesprávneho dátového typu počas výpočtu. Na to slúžia napríklad funkcie `numeric()`, `integer()`, `logical`, `character()`, ktorých argumentom je dĺžka vektora, ktoré napríklad v prípade dátového typu `logical` vytvoria vektor dĺžky desať s hodnotami `FALSE`, v prípade dátového typu `numeric` vytvoria vektor dĺžky desať s nulovými hodnotami alebo v prípade dátového typu `character` vytvoria vektor dĺžky desať s hodnotami prázdneho textového reťazca:

```

1 > initial.numeric <- numeric(10)
2 > initial.integer <- integer(10)
3 > initial.logical <- logical(10)
4 > initial.character <- character(10)
5 >
6 > typeof(initial.numeric)
7 [1] "double"
8 > typeof(initial.integer)
9 [1] "integer"
10 > typeof(initial.logical)
11 [1] "logical"
12 > typeof(initial.character)
13 [1] "character"
14 >
15 > print(initial.numeric)
16 [1] 0 0 0 0 0 0 0 0 0 0

```

```

17 > print(initial.integer)
18 [1] 0 0 0 0 0 0 0 0 0 0
19 > print(initial.logical)
20 [1] FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE
21 > print(initial.character)
22 [1] "" "" "" "" "" "" "" "" "" ""
23 >
24 > length(initial.numeric)
25 [1] 10
26 > length(initial.integer)
27 [1] 10
28 > length(initial.logical)
29 [1] 10
30 > length(initial.character)
31 [1] 10

```

6.1.2 Atomická matica

Atomická matica predstavuje základnú dátovú štruktúru v **R**. Ide o dvojrozmernú tabuľku, ktorej všetky prvky majú rovnaký dátový typ. Maticu si môžeme predstaviť ako vektory rovnakého dátového typu, ktoré sú spojené buď po riadkoch, alebo stĺpcoch do dvojrozmerného objektu. Nevýhodou matíc je práve homogenita ich prvkov. Pri reálnych analýzach máme niektoré premenné, ktoré sú celočíselné (vek), niektoré sú kategorické (pohlavie) a ak by sme chceli vstupné údaje analýzy reprezentovať atomickou maticou, podobne ako je to u vektorov, došlo by ku konverzii na nadradený dátový typ. Ak sme napríklad premennú pohlavie mali označené textovým reťazcom ("muž"/ "žena", "M"/ "F"), potom aj ostatné premenné (napríklad vek) budú prekonvertované na textový reťazec (dátový typ `character`). Práve z týchto dôvodov sa v praktických analýzach atomické matice veľmi nevyužívajú. Svoje použitie však majú pri realizácii lineárnej maticovej algebry (napríklad v regresii často využívaná) a taktiež niektoré funkcie na vstupe očakávajú alebo výstupe generujú atomické matice.

Funkcia `matrix()`

Atomickú maticu v **R** môžeme vytvoriť viacerými spôsobmi. Prvým a najviac používaným spôsobom je využitie funkcie `matrix()`, ktorej najdôležitejšie argumenty sú nasledovné:

- `data`: vektor hodnôt, z ktorých sa má matica vytvoriť,
- `nrow`: požadovaný počet riadkov matice,
- `ncol`: požadovaný počet stĺpcov matice,
- `byrow`: logická hodnota, prednastavená na `FALSE`, v prípade, ak je nastavená na `TRUE`, hodnoty sa do matice vyplňajú z vektora v argumente `data` po riadkoch, v prípade,

ak je nastavená na FALSE, hodnoty sa do matice vyplňajú z vektora v argumente `data` po stĺpcoch,

- **dimnames:** vo východiskovom tvare je nastavený na NULL, pomocou tohto argumentu je možné pomenovať dimenzie matice vo forme zoznamu (listu) (súhrnne riadky jedným názvom a súhrnne stĺpce jedným názvom: ako vytvárať menovky riadkov a stĺpcov je uvedené nižšie).

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > matrix(
5 +   data = vektor1,
6 +   nrow = 2,
7 +   ncol = 5,
8 +   byrow = FALSE
9 + )
10      [,1] [,2] [,3] [,4] [,5]
11 [1,]    1    3    5    7    9
12 [2,]    2    4    6    8   10
```

Všimnime si, že niektoré argumenty sú zbytočné, napríklad ak vytvárame maticu z vektora dĺžky desať a nadefinujeme počet riadkov na dva, nie je potrebné definovať počet stĺpcov, lebo ten sa automaticky dopočíta, keďže desiatka je deliteľná dvomi bezo zvyšku. Rovnako ak nedefinujeme počet stĺpcov, tak počet riadkov sa dopočíta automaticky. Argument `byrow` je vo východiskovom tvare nastavený na FALSE, pokiaľ chceme, aby sa hodnoty z vektora do matice vyplňali po stĺpcoch, nie je potrebné ho uvádzať:

```
1 > matrix(data = vektor1,
2 +       nrow = 2,)
3      [,1] [,2] [,3] [,4] [,5]
4 [1,]    1    3    5    7    9
5 [2,]    2    4    6    8   10
6 > matrix(data = vektor1,
7 +       ncol = 5,)
8      [,1] [,2] [,3] [,4] [,5]
9 [1,]    1    3    5    7    9
10 [2,]    2    4    6    8   10
```

Ak však chceme, aby sa hodnoty z vektora do matice vyplňali po riadkoch, je potrebné uviesť argument `byrow = TRUE`:

```
1 > matrix(data = vektor1,
2 +       ncol = 5,)
3      [,1] [,2] [,3] [,4] [,5]
4 [1,]    1    3    5    7    9
```

```

5 [2,]    2    4    6    8   10
6 >
7 > matrix(data = vektor1,
8 +       ncol = 5,
9 +       byrow = TRUE)
10      [,1] [,2] [,3] [,4] [,5]
11 [1,]    1    2    3    4    5
12 [2,]    6    7    8    9   10

```

Podobne, ako to bolo v prípade vektorov, aj pri tvorbe matíc sa **R** snaží „recyklovať“ dáta. Ak má napríklad vektor, z ktorého vytvárame maticu dĺžku menšiu ako je súčin počtu riadkov a počtu stĺpcov matice, ktorú chceme vytvoriť, **R** začne chýbajúce hodnoty vyplňať od začiatku vektora. To býva často zdrojom chýb, práve preto pri každej takejto recyklácii nasleduje upozornenie, aby sme si uvedomili, že z vektora, ktorý uvažujeme, nie je možné vytvoriť maticu daného rozmeru bez toho, aby sa jeho hodnoty nezačali opakovať. Uvažujme napríklad, že chceme vytvoriť maticu rozmeru 3×4 z vektora uloženého v objekte **vektor1**, o ktorom vieme, že nemá požadovanú dĺžku 12, ale iba 10:

```

1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > length(vektor1)
5 [1] 10
6 > matrix(data = vektor1,
7 +       nrow = 3,
8 +       ncol = 4)
9      [,1] [,2] [,3] [,4]
10 [1,]    1    4    7    10
11 [2,]    2    5    8     1
12 [3,]    3    6    9     2
13 Warning message:
14 In matrix(data = vektor1, nrow = 3, ncol = 4) :
15   data length [10] is not a sub-multiple or multiple of the number
    of rows [3]

```

Vidíme, že sa vytvorila matica rozmeru 3×4 z vektora uloženého v objekte **vektor1**, pričom hodnoty sa vyplňali po stĺpcoch (východiskové nastavenie). Hodnota v prvom riadku a štvrtom stĺpci predstavuje poslednú hodnotu vektora uloženého v objekte **vektor1**. Hodnota v druhom riadku a štvrtom stĺpci by za normálnych okolností vyžadovala jedenástu hodnotu vektora, avšak keďže tento vektor je len dĺžky desať, na túto pozíciu sa uložila prvá hodnota vektora. Hodnota v treťom riadku a štvrtom stĺpci by za normálnych okolností vyžadovala dvanástu hodnotu vektora, avšak keďže tento vektor je len dĺžky desať, na túto pozíciu sa uložila druhá hodnota vektora. Na toto vyplňovanie od začiatku sme boli upozornení textovým výstupom.

Funkcie rbind() a cbind()

Matice môžeme vytvoriť aj skladaním viacerých vektorov po riadkoch (**r**: row) alebo po stĺpcoch (**c**: column), na čo využívame funkcie **rbind()** a **cbind()**. Pomocou týchto funkcií môžeme skladať aj samotné matice. Ak by vektory nemali rovnakú dĺžku, podobne ako v prípade funkcie **matrix()**, dochádza k „recyklácii“ hodnôt, na čo nás **R** opäť upozorní.

```
1 > print(konverze1)
2 [1] 1 0 1 1 5
3 > length(konverze1)
4 [1] 5
5 > print(vektor5)
6 [1] 2 9 16 23
7 > length(vektor5)
8 [1] 4
9 > rbind(konverze1, vektor5)
10      [,1] [,2] [,3] [,4] [,5]
11 konverze1    1    0    1    1    5
12 vektor5      2    9   16   23    2
13 Warning message:
14 In rbind(konverze1, vektor5) :
15   number of columns of result is not a multiple of vector length (
16     arg 2)
17 > cbind(konverze1, vektor5)
18      konverze1 vektor5
19 [1,]          1          2
20 [2,]          0          9
21 [3,]          1         16
22 [4,]          1         23
23 [5,]          5          2
24 Warning message:
25 In cbind(konverze1, vektor5) :
26   number of rows of result is not a multiple of vector length (arg
27     2)
```

Všimnime si, že pri aplikácii funkcie **cbind()** sa automaticky stĺpce matice pomenujú na základe názvov vektorov, z ktorých sa matica vytvára, a pri aplikácii funkcie **rbind()** sa automaticky riadky matice pomenujú na základe vektorov, z ktorých sa matica vytvára.

Vytvorme si vektory **a**, **b**, **c**, **d** dĺžky šesť náhodných celých čísel od 1 do 100, pričom hodnoty v jednotlivých vektoroch sa nemôžu opakovať a vytvorme z týchto vektorov maticu **mat.R** spojením týchto vektorov po riadkoch a maticu **mat.C** spojením týchto vektorov po stĺpcoch:

```
1 > set.seed(123)
2 > a <- sample(1:100, 6, replace = FALSE)
```

```

3 > print(a)
4 [1] 31 79 51 14 67 42
5 > b <- sample(1:100, 6, replace = FALSE)
6 > print(b)
7 [1] 50 43 14 25 90 91
8 > c <- sample(1:100, 6, replace = FALSE)
9 > print(c)
10 [1] 69 91 57 92 9 93
11 > d <- sample(1:100, 6, replace = FALSE)
12 > print(d)
13 [1] 99 72 26 7 42 9
14 > mat.R <- rbind(a, b, c, d)
15 > mat.C <- cbind(a, b, c, d)
16 > print(mat.R)
17      [,1] [,2] [,3] [,4] [,5] [,6]
18 a      31   79   51   14   67   42
19 b      50   43   14   25   90   91
20 c      69   91   57   92    9   93
21 d      99   72   26    7   42    9
22 > print(mat.C)
23           a  b  c  d
24 [1,]  31 50 69 99
25 [2,]  79 43 91 72
26 [3,]  51 14 57 26
27 [4,]  14 25 92  7
28 [5,]  67 90  9 42
29 [6,]  42 91 93  9

```

Ak chceme zistiť rozmer matice, môžeme použiť funkciu `dim()`, ktorej argumentom je objekt, ktorého rozmery chceme poznať. V prípade, že ide o maticu, tak na výstupe je uvedený vektor s dvoma hodnotami, prvá hodnota predstavuje počet riadkov a druhá hodnota predstavuje počet stĺpcov matice.

Ak chceme zistiť počet riadkov matice, môžeme použiť funkciu `nrow()`, ktorej argumentom je objekt, predstavujúci vektor, maticu, pole alebo dataset. V prípade, že ide o maticu alebo dataset, tak na výstupe je uvedená hodnota predstavujúca počet riadkov tejto matice, respektíve datasetu. Alternatívou je z vektora, ktorý vyprodukuje funkcia `dim()`, vytiahnuť hodnotu, nachádzajúcu sa na prvom mieste.

Ak chceme zistiť počet stĺpcov matice, môžeme použiť funkciu `ncol()`, ktorej argumentom je objekt, predstavujúci vektor, maticu, pole alebo dataset. V prípade, že ide o maticu alebo dataset, tak na výstupe je uvedená hodnota predstavujúca počet stĺpcov tejto matice, respektíve datasetu. Alternatívou je z vektora, ktorý vyprodukuje funkcia `dim()`, vytiahnuť hodnotu nachádzajúcu sa na druhom mieste.

Ak chceme zistiť, koľko hodnôt obsahuje matica (ide o súčin počtu riadkov a počtu stĺpcov

matice), využívame funkciu `length()`, ktorú poznáme už z práce s vektormi a ktorú sme používali, keď sme chceli poznať dĺžku (počet prvkov) vektora.

```
1 > print(mat.R)
2   [,1] [,2] [,3] [,4] [,5] [,6]
3 a    31    79    51    14    67    42
4 b    50    43    14    25    90    91
5 c    69    91    57    92     9    93
6 d    99    72    26     7    42     9
7 > dim(mat.R)
8 [1] 4 6
9 > nrow(mat.R)
10 [1] 4
11 > dim(mat.R)[1]
12 [1] 4
13 > ncol(mat.R)
14 [1] 6
15 > dim(mat.R)[2]
16 [1] 6
17 > length(mat.R)
18 [1] 24
19 > nrow(mat.R) * ncol(mat.R)
20 [1] 24
```

Vytvoríme inú maticu (ozn. `mat.RR`) so šiestimi stĺpcami a ukážeme si použitie funkcie `rbind()` na celých maticiach:

```
1 > set.seed(123)
2 > mat.RR <- matrix(runif(54), ncol = 6)
3 > print(mat.RR)
4           [,1]      [,2]      [,3]      [,4]      [,5]      [,6]
5 [1,] 0.2875775 0.45661474 0.3279207 0.59414202 0.7584595 0.13880606
6 [2,] 0.7883051 0.95683335 0.9545036 0.28915974 0.2164079 0.23303410
7 [3,] 0.4089769 0.45333416 0.8895393 0.14711365 0.3181810 0.46596245
8 [4,] 0.8830174 0.67757064 0.6928034 0.96302423 0.2316258 0.26597264
9 [5,] 0.9404673 0.57263340 0.6405068 0.90229905 0.1428000 0.85782772
10 [6,] 0.0455565 0.10292468 0.9942698 0.69070528 0.4145463 0.04583117
11 [7,] 0.5281055 0.89982497 0.6557058 0.79546742 0.4137243 0.44220007
12 [8,] 0.8924190 0.24608773 0.7085305 0.02461368 0.3688455 0.79892485
13 [9,] 0.5514350 0.04205953 0.5440660 0.47779597 0.1524447 0.12189926
```

Keďže majú matice `mat.R` a `mat.RR` rovnaký počet stĺpcov, môžeme na nich aplikovať spájanie po riadkoch (funkcia `rbind()`). Vzhľadom na to, že tieto matice nemajú rovnaký počet riadkov, nie je možné na nich aplikovať spájanie po stĺpcoch (funkcia `cbind()`).

```
1 > ncol(mat.R)
2 [1] 6
```

```

3 > ncol(mat.RR)
4 [1] 6
5 > rbind(mat.R, mat.RR)
6      [,1]      [,2]      [,3]      [,4]      [,5]
7 a 31.0000000 79.0000000 51.0000000 14.0000000 67.0000000
8 b 50.0000000 43.0000000 14.0000000 25.0000000 90.0000000
9 c 69.0000000 91.0000000 57.0000000 92.0000000  9.0000000
10 d 99.0000000 72.0000000 26.0000000  7.0000000 42.0000000
11    0.2875775  0.45661474  0.3279207  0.59414202  0.7584595
12    0.7883051  0.95683335  0.9545036  0.28915974  0.2164079
13    0.4089769  0.45333416  0.8895393  0.14711365  0.3181810
14    0.8830174  0.67757064  0.6928034  0.96302423  0.2316258
15    0.9404673  0.57263340  0.6405068  0.90229905  0.1428000
16    0.0455565  0.10292468  0.9942698  0.69070528  0.4145463
17    0.5281055  0.89982497  0.6557058  0.79546742  0.4137243
18    0.8924190  0.24608773  0.7085305  0.02461368  0.3688455
19    0.5514350  0.04205953  0.5440660  0.47779597  0.1524447
20      [,6]
21 a 42.00000000
22 b 91.00000000
23 c 93.00000000
24 d  9.00000000
25    0.13880606
26    0.23303410
27    0.46596245
28    0.26597264
29    0.85782772
30    0.04583117
31    0.44220007
32    0.79892485
33    0.12189926
34 > nrow(mat.R)
35 [1] 4
36 > nrow(mat.RR)
37 [1] 9
38 > cbind(mat.R, mat.RR)
39 Error in cbind(mat.R, mat.RR) :
40   number of rows of matrices must match (see arg 2)

```

Uložme si po riadkoch spojenú maticu do objektu `mat.RRR` a pozrime sa niektoré jej vlastnosti:

```

1 > mat.RRR <- rbind(mat.R, mat.RR)
2 > mat.RRR
3      [,1]      [,2]      [,3]      [,4]      [,5]
4 a 31.0000000 79.0000000 51.0000000 14.0000000 67.0000000

```

```

5 b 50.0000000 43.0000000 14.0000000 25.0000000 90.0000000
6 c 69.0000000 91.0000000 57.0000000 92.0000000 9.0000000
7 d 99.0000000 72.0000000 26.0000000 7.0000000 42.0000000
8 0.2875775 0.45661474 0.3279207 0.59414202 0.7584595
9 0.7883051 0.95683335 0.9545036 0.28915974 0.2164079
10 0.4089769 0.45333416 0.8895393 0.14711365 0.3181810
11 0.8830174 0.67757064 0.6928034 0.96302423 0.2316258
12 0.9404673 0.57263340 0.6405068 0.90229905 0.1428000
13 0.0455565 0.10292468 0.9942698 0.69070528 0.4145463
14 0.5281055 0.89982497 0.6557058 0.79546742 0.4137243
15 0.8924190 0.24608773 0.7085305 0.02461368 0.3688455
16 0.5514350 0.04205953 0.5440660 0.47779597 0.1524447
17 [,6]
18 a 42.00000000
19 b 91.00000000
20 c 93.00000000
21 d 9.00000000
22 0.13880606
23 0.23303410
24 0.46596245
25 0.26597264
26 0.85782772
27 0.04583117
28 0.44220007
29 0.79892485
30 0.12189926
31 > dim(mat.RRR)
32 [1] 13 6
33 > length(mat.RRR)
34 [1] 78

```

Vidíme, že táto matica má pomenované iba prvé štyri riadky a nemá pomenované stĺpce. Pomenujme riadky tejto matice malými písmenami latinskej abecedy a stĺpce tejto matice veľkými písmenami latinskej abecedy. Na to slúžia funkcie `rownames()` a `colnames()`. Ich argumentom je matica alebo dataset a môžeme nimi jednak nechať tieto názvy vypísať a taktiež vieme tieto názvy pomocou týchto funkcií modifikovať. Naša matica má rozmery 13×6 :

```

1 > dim(mat.RRR)
2 [1] 13 6
3 > length(mat.RRR)
4 [1] 78
5 > dim(mat.RRR)
6 [1] 13 6
7 > rownames(mat.RRR)
8 [1] "a" "b" "c" "d" "" "" "" "" "" "" "" "" ""

```

```

9 > colnames(mat.RRR)
10 NULL
11 > rownames(mat.RRR) <- letters[1:13]
12 > colnames(mat.RRR) <- LETTERS[1:6]
13 > mat.RRR
14           A           B           C           D           E
15 a 31.0000000 79.0000000 51.0000000 14.0000000 67.0000000
16 b 50.0000000 43.0000000 14.0000000 25.0000000 90.0000000
17 c 69.0000000 91.0000000 57.0000000 92.0000000  9.0000000
18 d 99.0000000 72.0000000 26.0000000  7.0000000 42.0000000
19 e  0.2875775  0.45661474  0.3279207  0.59414202  0.7584595
20 f  0.7883051  0.95683335  0.9545036  0.28915974  0.2164079
21 g  0.4089769  0.45333416  0.8895393  0.14711365  0.3181810
22 h  0.8830174  0.67757064  0.6928034  0.96302423  0.2316258
23 i  0.9404673  0.57263340  0.6405068  0.90229905  0.1428000
24 j  0.0455565  0.10292468  0.9942698  0.69070528  0.4145463
25 k  0.5281055  0.89982497  0.6557058  0.79546742  0.4137243
26 l  0.8924190  0.24608773  0.7085305  0.02461368  0.3688455
27 m  0.5514350  0.04205953  0.5440660  0.47779597  0.1524447
28           F
29 a 42.00000000
30 b 91.00000000
31 c 93.00000000
32 d  9.00000000
33 e  0.13880606
34 f  0.23303410
35 g  0.46596245
36 h  0.26597264
37 i  0.85782772
38 j  0.04583117
39 k  0.44220007
40 l  0.79892485
41 m  0.12189926

```

Aby sa nám s hodnotami matice lepšie pracovalo, môžeme si ich zaokrúhliť napríklad na dve desatinné miesta. Využívame na to už známu funkciu `round()`:

```

1 > mat.RRR <- round(mat.RRR, digits = 2)
2 > mat.RRR
3           A           B           C           D           E           F
4 a 31.00 79.00 51.00 14.00 67.00 42.00
5 b 50.00 43.00 14.00 25.00 90.00 91.00
6 c 69.00 91.00 57.00 92.00  9.00 93.00
7 d 99.00 72.00 26.00  7.00 42.00  9.00
8 e  0.29  0.46  0.33  0.59  0.76  0.14
9 f  0.79  0.96  0.95  0.29  0.22  0.23

```

10	g	0.41	0.45	0.89	0.15	0.32	0.47
11	h	0.88	0.68	0.69	0.96	0.23	0.27
12	i	0.94	0.57	0.64	0.90	0.14	0.86
13	j	0.05	0.10	0.99	0.69	0.41	0.05
14	k	0.53	0.90	0.66	0.80	0.41	0.44
15	l	0.89	0.25	0.71	0.02	0.37	0.80
16	m	0.55	0.04	0.54	0.48	0.15	0.12

Ak chceme pracovať s hodnotami z konkrétného riadka, respektíve stĺpca, využívame rovnako ako v prípade vektora hranaté zátvorky. Keďže matica je dvojrozmerný objekt, prvá hodnota v hranatej zátvorke predstavuje poradie riadka a druhá hodnota v hranatej zátvorke predstavuje poradie stĺpca.

Chceme napríklad hodnotu v druhom riadku (b) a v štvrtom stĺpci (D):

```
1 > mat.RRR[2, 4]
2 [1] 25
```

Ak máme pomenované riadky matice, ako v tomto prípade, je možný aj prístup cez menovky riadkov a stĺpcov, čo však môže byť častým zdrojom chýb, preto to veľmi neodporúčame.

```
1 > mat.RRR["b", "D"]
2 [1] 25
```

Ak chceme všetky hodnoty v treťom riadku (c), využijeme predchádzajúcu notáciu, pri ktorej ako prvú dimenziu označíme poradie riadka a poradie stĺpca necháme prázdne, čím vlastne vyberáme hodnoty z konkrétného riadka a všetkých stĺpcov. Opäť môžeme použiť aj prístup cez menovky riadkov.

```
1 > mat.RRR[3, ]
2   A   B   C   D   E   F
3 69 91 57 92   9 93
4 > mat.RRR["c", ]
5   A   B   C   D   E   F
6 69 91 57 92   9 93
```

Úplne rovnaký princíp platí aj pri vyberaní hodnôt z konkrétného stĺpca, povedzme, že chceme vybrať všetky hodnoty z piateho stĺpca (E):

```
1 > mat.RRR[, 5]
2      a      b      c      d      e      f      g      h      i      j      k
3 67.00 90.00  9.00 42.00  0.76  0.22  0.32  0.23  0.14  0.41  0.41
4      l      m
5  0.37  0.15
6 > mat.RRR[, "E"]
7      a      b      c      d      e      f      g      h      i      j      k
```

```

8 67.00 90.00 9.00 42.00 0.76 0.22 0.32 0.23 0.14 0.41 0.41
9      1      m
10 0.37 0.15

```

Ak chceme vybrať napríklad viaceré riadky, je potrebné ich poradie uviesť vo vektore. V prípade, ak ide o za sebou idúce riadky, môžeme jednoducho použiť operátor aritmetickej postupnosti s diferenciou 1, respektíve -1 (:). Povedzme, že chceme zobrazíť všetky hodnoty v štvrtom, piatom a šiestom riadku (d, e, f):

```

1 > mat.RRR[4:6, ]
2      A      B      C      D      E      F
3 d 99.00 72.00 26.00 7.00 42.00 9.00
4 e  0.29  0.46  0.33 0.59  0.76 0.14
5 f  0.79  0.96  0.95 0.29  0.22 0.23
6 > mat.RRR[c(4, 5, 6), ]
7      A      B      C      D      E      F
8 d 99.00 72.00 26.00 7.00 42.00 9.00
9 e  0.29  0.46  0.33 0.59  0.76 0.14
10 f  0.79  0.96  0.95 0.29  0.22 0.23
11 > mat.RRR[c("d", "e", "f"), ]
12      A      B      C      D      E      F
13 d 99.00 72.00 26.00 7.00 42.00 9.00
14 e  0.29  0.46  0.33 0.59  0.76 0.14
15 f  0.79  0.96  0.95 0.29  0.22 0.23

```

Ak chceme vybrať viaceré stĺpce, taktiež je potrebné ich poradie uviesť vo vektore. Na príklade stĺpcov si ukážeme, čo robiť, ak ide o stĺpce, ktoré nenasledujú za sebou. Povedzme, že chceme všetky hodnoty v prvom, treťom a šiestom stĺpci (A, C, F):

```

1 > mat.RRR[, c(1, 3, 6)]
2      A      C      F
3 a 31.00 51.00 42.00
4 b 50.00 14.00 91.00
5 c 69.00 57.00 93.00
6 d 99.00 26.00 9.00
7 e  0.29  0.33 0.14
8 f  0.79  0.95 0.23
9 g  0.41  0.89 0.47
10 h  0.88  0.69 0.27
11 i  0.94  0.64 0.86
12 j  0.05  0.99 0.05
13 k  0.53  0.66 0.44
14 l  0.89  0.71 0.80
15 m  0.55  0.54 0.12
16 > mat.RRR[, c("A", "C", "F")]
17      A      C      F

```

```

18 a 31.00 51.00 42.00
19 b 50.00 14.00 91.00
20 c 69.00 57.00 93.00
21 d 99.00 26.00 9.00
22 e 0.29 0.33 0.14
23 f 0.79 0.95 0.23
24 g 0.41 0.89 0.47
25 h 0.88 0.69 0.27
26 i 0.94 0.64 0.86
27 j 0.05 0.99 0.05
28 k 0.53 0.66 0.44
29 l 0.89 0.71 0.80
30 m 0.55 0.54 0.12

```

Rovnako ako pri vektoroch aj pri maticiach môžeme vybrať tie riadky, respektíve stĺpce, s ktorými nechceme pracovať, napríklad si zobrazíme všetky okrem tretieho a siedmeho riadka matice v objekte `mat.RRR` a všetko okrem posledných dvoch stĺpcov:

```

1 > print(mat.RRR)
2      A      B      C      D      E      F
3 a 31.00 79.00 51.00 14.00 67.00 42.00
4 b 50.00 43.00 14.00 25.00 90.00 91.00
5 c 69.00 91.00 57.00 92.00 9.00 93.00
6 d 99.00 72.00 26.00 7.00 42.00 9.00
7 e 0.29 0.46 0.33 0.59 0.76 0.14
8 f 0.79 0.96 0.95 0.29 0.22 0.23
9 g 0.41 0.45 0.89 0.15 0.32 0.47
10 h 0.88 0.68 0.69 0.96 0.23 0.27
11 i 0.94 0.57 0.64 0.90 0.14 0.86
12 j 0.05 0.10 0.99 0.69 0.41 0.05
13 k 0.53 0.90 0.66 0.80 0.41 0.44
14 l 0.89 0.25 0.71 0.02 0.37 0.80
15 m 0.55 0.04 0.54 0.48 0.15 0.12
16 > dim(mat.RRR)
17 [1] 13 6
18 > mat.RRR[-c(3, 7), -(5:6)]
19      A      B      C      D
20 a 31.00 79.00 51.00 14.00
21 b 50.00 43.00 14.00 25.00
22 d 99.00 72.00 26.00 7.00
23 e 0.29 0.46 0.33 0.59
24 f 0.79 0.96 0.95 0.29
25 h 0.88 0.68 0.69 0.96
26 i 0.94 0.57 0.64 0.90
27 j 0.05 0.10 0.99 0.69
28 k 0.53 0.90 0.66 0.80

```

```

29 1  0.89  0.25  0.71  0.02
30 m  0.55  0.04  0.54  0.48

```

Práve pri práci s maticami a datasetmi, ktoré majú množstvo pozorovaní (riadkov), je veľmi vhodné použitie funkcií `head()`, `tail()` a `summary()`, ktoré sme predstavovali už pri práci s vektormi. Funkcia `head()` v prípade, že na vstupe má maticu alebo dataset, na výstupe vypíše hodnoty prvých pár (prednastavené na šesť, možno zmeniť argumentom `n`) riadkov a im odpovedajúcich stĺpcov, aby sme si vedeli predstaviť, s akými údajmi pracujeme. Funkcia `tail()` funguje úplne rovnako, akurát, že nevypisuje hodnoty v prvých, ale v posledných riadkoch. Funkcia `summary()` vypisuje základné popisné štatistiky hodnôt po stĺpcoch. Vzhľadom na to, že matica je homogénnou dátovou štruktúrou a všetky jej hodnoty sú rovnakého dátového typu, mohli by sme požadovať aj výpočet popisných štatistík po riadkoch. Na to by sme museli maticu transponovať (vymeniť riadky za stĺpce), na čo slúži funkcia `t()`. Ak by sme túto funkciu aplikovali na vektor, vedeli by sme riadkový vektor zmeniť na stĺpcový a naopak, čo je potrebné napríklad pri výpočtoch v lineárnej algebre, ktoré sa využívajú v regresii a podobne.

```

1  > print(mat.RRR)
2      A      B      C      D      E      F
3 a 31.00 79.00 51.00 14.00 67.00 42.00
4 b 50.00 43.00 14.00 25.00 90.00 91.00
5 c 69.00 91.00 57.00 92.00  9.00 93.00
6 d 99.00 72.00 26.00  7.00 42.00  9.00
7 e  0.29  0.46  0.33  0.59  0.76  0.14
8 f  0.79  0.96  0.95  0.29  0.22  0.23
9 g  0.41  0.45  0.89  0.15  0.32  0.47
10 h  0.88  0.68  0.69  0.96  0.23  0.27
11 i  0.94  0.57  0.64  0.90  0.14  0.86
12 j  0.05  0.10  0.99  0.69  0.41  0.05
13 k  0.53  0.90  0.66  0.80  0.41  0.44
14 l  0.89  0.25  0.71  0.02  0.37  0.80
15 m  0.55  0.04  0.54  0.48  0.15  0.12
16 > head(mat.RRR)
17      A      B      C      D      E      F
18 a 31.00 79.00 51.00 14.00 67.00 42.00
19 b 50.00 43.00 14.00 25.00 90.00 91.00
20 c 69.00 91.00 57.00 92.00  9.00 93.00
21 d 99.00 72.00 26.00  7.00 42.00  9.00
22 e  0.29  0.46  0.33  0.59  0.76  0.14
23 f  0.79  0.96  0.95  0.29  0.22  0.23
24 > tail(mat.RRR)
25      A      B      C      D      E      F
26 h 0.88 0.68 0.69 0.96 0.23 0.27
27 i 0.94 0.57 0.64 0.90 0.14 0.86
28 j 0.05 0.10 0.99 0.69 0.41 0.05

```



```

29 k 0.53 0.90 0.66 0.80 0.41 0.44
30 l 0.89 0.25 0.71 0.02 0.37 0.80
31 m 0.55 0.04 0.54 0.48 0.15 0.12
32 > summary(mat.RRR)
33           A           B           C           D
34 Min.      : 0.05      Min.      : 0.04      Min.      : 0.33      Min.      : 0.02
35 1st Qu.: 0.53      1st Qu.: 0.45      1st Qu.: 0.66      1st Qu.: 0.48
36 Median : 0.88      Median : 0.68      Median : 0.89      Median : 0.80
37 Mean     :19.56      Mean     :22.26      Mean     :11.88      Mean     :10.99
38 3rd Qu.:31.00      3rd Qu.:43.00      3rd Qu.:14.00      3rd Qu.: 7.00
39 Max.     :99.00      Max.     :91.00      Max.     :57.00      Max.     :92.00
40           E           F
41 Min.      : 0.14      Min.      : 0.05
42 1st Qu.: 0.23      1st Qu.: 0.23
43 Median : 0.41      Median : 0.47
44 Mean     :16.23      Mean     :18.34
45 3rd Qu.: 9.00      3rd Qu.: 9.00
46 Max.     :90.00      Max.     :93.00
47 > summary(t(mat.RRR))
48           a           b           c           d
49 Min.      :14.00      Min.      :14.00      Min.      : 9.00      Min.      : 7.00
50 1st Qu.:33.75      1st Qu.:29.50      1st Qu.:60.00      1st Qu.:13.25
51 Median :46.50      Median :46.50      Median :80.00      Median :34.00
52 Mean     :47.33      Mean     :52.17      Mean     :68.50      Mean     :42.50
53 3rd Qu.:63.00      3rd Qu.:80.00      3rd Qu.:91.75      3rd Qu.:64.50
54 Max.     :79.00      Max.     :91.00      Max.     :93.00      Max.     :99.00
55           e           f           g
56 Min.      :0.1400      Min.      :0.2200      Min.      :0.1500
57 1st Qu.:0.3000      1st Qu.:0.2450      1st Qu.:0.3425
58 Median :0.3950      Median :0.5400      Median :0.4300
59 Mean     :0.4283      Mean     :0.5733      Mean     :0.4483
60 3rd Qu.:0.5575      3rd Qu.:0.9100      3rd Qu.:0.4650
61 Max.     :0.7600      Max.     :0.9600      Max.     :0.8900
62           h           i           j
63 Min.      :0.2300      Min.      :0.1400      Min.      :0.0500
64 1st Qu.:0.3725      1st Qu.:0.5875      1st Qu.:0.0625
65 Median :0.6850      Median :0.7500      Median :0.2550
66 Mean     :0.6183      Mean     :0.6750      Mean     :0.3817
67 3rd Qu.:0.8325      3rd Qu.:0.8900      3rd Qu.:0.6200
68 Max.     :0.9600      Max.     :0.9400      Max.     :0.9900
69           k           l           m
70 Min.      :0.4100      Min.      :0.0200      Min.      :0.0400
71 1st Qu.:0.4625      1st Qu.:0.2800      1st Qu.:0.1275
72 Median :0.5950      Median :0.5400      Median :0.3150
73 Mean     :0.6233      Mean     :0.5067      Mean     :0.3133
74 3rd Qu.:0.7650      3rd Qu.:0.7775      3rd Qu.:0.5250

```

75	Max.	: 0.9000	Max.	: 0.8900	Max.	: 0.5500
----	------	----------	------	----------	------	----------

Atribúty matice

Prostredníctvom funkcie `attributes()`, ktorej argumentom je ľubovoľný objekt, si môžeme vypísať základné atribúty tohto objektu (rozmer, prípadne názvy riadkov alebo názvy stĺpcov). Ilustrujme si použitie funkcie na vektore bez menoviek (objekt `a`), vektore s menovkami (objekt `vektor1`), matici iba s menovkami stĺpcov (objekt `mat.C`) a matici aj s menovkami riadkov aj s menovkami stĺpcov (objekt `mat.RRR`):

```
1 > print(a)
2 [1] 31 79 51 14 67 42
3 > typeof(a)
4 [1] "integer"
5 > attributes(a)
6 NULL
7 >
8 > print(vektor1)
9  A  B  C  D  E  F  G  H  I  J
10 1  2  3  4  5  6  7  8  9 10
11 > typeof(vektor1)
12 [1] "integer"
13 > attributes(vektor1)
14 $names
15 [1] "A" "B" "C" "D" "E" "F" "G" "H" "I" "J"
16
17 >
18 > print(mat.C)
19      a  b  c  d
20 [1,] 31 50 69 99
21 [2,] 79 43 91 72
22 [3,] 51 14 57 26
23 [4,] 14 25 92  7
24 [5,] 67 90  9 42
25 [6,] 42 91 93  9
26 > typeof(mat.C)
27 [1] "integer"
28 > attributes(mat.C)
29 $dim
30 [1] 6 4
31
32 $dimnames
33 $dimnames[[1]]
34 NULL
35
```

```

36 $dimnames[[2]]
37 [1] "a" "b" "c" "d"
38
39
40 >
41 > print(mat.RRR)
42      A      B      C      D      E      F
43 a 31.00 79.00 51.00 14.00 67.00 42.00
44 b 50.00 43.00 14.00 25.00 90.00 91.00
45 c 69.00 91.00 57.00 92.00  9.00 93.00
46 d 99.00 72.00 26.00  7.00 42.00  9.00
47 e  0.29  0.46  0.33  0.59  0.76  0.14
48 f  0.79  0.96  0.95  0.29  0.22  0.23
49 g  0.41  0.45  0.89  0.15  0.32  0.47
50 h  0.88  0.68  0.69  0.96  0.23  0.27
51 i  0.94  0.57  0.64  0.90  0.14  0.86
52 j  0.05  0.10  0.99  0.69  0.41  0.05
53 k  0.53  0.90  0.66  0.80  0.41  0.44
54 l  0.89  0.25  0.71  0.02  0.37  0.80
55 m  0.55  0.04  0.54  0.48  0.15  0.12
56 > typeof(mat.RRR)
57 [1] "double"
58 > attributes(mat.RRR)
59 $dim
60 [1] 13  6
61
62 $dimnames
63 $dimnames[[1]]
64 [1] "a" "b" "c" "d" "e" "f" "g" "h" "i" "j" "k" "l" "m"
65
66 $dimnames[[2]]
67 [1] "A" "B" "C" "D" "E" "F"

```

Všimnime si, že ak si odmyslíme názvy riadkov a stĺpcov, tak atomická matica je vlastne iba atomickým vektorom, ktorý má navyše definovaný atribút `dim` predstavujúci jej rozmer (počet riadkov a stĺpcov). Preto vytvorenie matice môže vyzeráť aj tak, že ľubovoľnému atomickému vektoru nadefinujeme atribút `dim`, napríklad pomocou funkcie `dim()` alebo priamo nadefinovaním príslušného atribútu vo funkcii `attributes()` s využitím operátora `$`:

```

1 > postupnost1 <- 1:36
2 > print(postupnost1)
3 [1]  1  2  3  4  5  6  7  8  9 10 11 12 13 14 15 16 17 18 19 20 21
4 [22] 22 23 24 25 26 27 28 29 30 31 32 33 34 35 36
5 > length(postupnost1)
6 [1] 36

```

```

7 > typeof(postupnost1)
8 [1] "integer"
9 > attributes(postupnost1)
10 NULL
11 > dim(postupnost1) <- c(6, 6)
12 > print(postupnost1)
13      [,1] [,2] [,3] [,4] [,5] [,6]
14 [1,]     1     7    13    19    25    31
15 [2,]     2     8    14    20    26    32
16 [3,]     3     9    15    21    27    33
17 [4,]     4    10    16    22    28    34
18 [5,]     5    11    17    23    29    35
19 [6,]     6    12    18    24    30    36
20 > attributes(postupnost1)
21 $dim
22 [1] 6 6
23
24 >
25 > postupnost2 <- 1:25
26 > print(postupnost2)
27 [1]  1  2  3  4  5  6  7  8  9 10 11 12 13 14 15 16 17 18 19 20 21
28 [22] 22 23 24 25
29 > length(postupnost2)
30 [1] 25
31 > typeof(postupnost2)
32 [1] "integer"
33 > attributes(postupnost2)
34 NULL
35 > attributes(postupnost2)$dim <- c(5, 5)
36 > print(postupnost2)
37      [,1] [,2] [,3] [,4] [,5]
38 [1,]     1     6    11    16    21
39 [2,]     2     7    12    17    22
40 [3,]     3     8    13    18    23
41 [4,]     4     9    14    19    24
42 [5,]     5    10    15    20    25
43 > attributes(postupnost2)
44 $dim
45 [1] 5 5

```

To isté platí aj v prípade, ak by sme chceli atomickú maticu zmeniť späť na atomický vektor, vtedy postačí atribút dim nastaviť na NULL.

```

1 > print(mat.C)
2      a  b  c  d
3 [1,] 31 50 69 99

```

```

4 [2,] 79 43 91 72
5 [3,] 51 14 57 26
6 [4,] 14 25 92 7
7 [5,] 67 90 9 42
8 [6,] 42 91 93 9
9 > typeof(mat.C)
10 [1] "integer"
11 > attributes(mat.C)
12 $dim
13 [1] 6 4
14
15 $dimnames
16 $dimnames[[1]]
17 NULL
18
19 $dimnames[[2]]
20 [1] "a" "b" "c" "d"
21
22
23 > dim(mat.C)<-NULL
24 > print(mat.C)
25 [1] 31 79 51 14 67 42 50 43 14 25 90 91 69 91 57 92 9 93 99 72 26
26 [22] 7 42 9

```

Podobne funguje aj funkcia `as.vector()`:

```

1 > print(mat.R)
2 [,1] [,2] [,3] [,4] [,5] [,6]
3 a 74 93 75 44 61 54
4 b 66 33 21 19 79 45
5 c 6 67 90 36 74 99
6 d 4 38 79 39 66 88
7 > typeof(mat.R)
8 [1] "integer"
9 > attributes(mat.R)
10 $dim
11 [1] 4 6
12
13 $dimnames
14 $dimnames[[1]]
15 [1] "a" "b" "c" "d"
16
17 $dimnames[[2]]
18 NULL
19
20 > as.vector(mat.R)

```

```

21 [1] 74 66 6 4 93 33 67 38 75 21 90 79 44 19 36 39 61 79 74 66
22 [21] 54 45 99 88

```

Na overenie, či ide o maticu, môžeme použiť funkciu `is.matrix()`. Podobne, ako pri vektorech, si atomickú štruktúru matice môžeme overiť prostredníctvom funkcie `is.atomic()`:

```

1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > is.matrix(vektor1)
5 [1] FALSE
6 > is.atomic(vektor1)
7 [1] TRUE
8 >
9 > print(mat.RRR)
10      A      B      C      D      E      F
11 a 31.00 79.00 51.00 14.00 67.00 42.00
12 b 50.00 43.00 14.00 25.00 90.00 91.00
13 c 69.00 91.00 57.00 92.00  9.00 93.00
14 d 99.00 72.00 26.00  7.00 42.00  9.00
15 e  0.29  0.46  0.33  0.59  0.76  0.14
16 f  0.79  0.96  0.95  0.29  0.22  0.23
17 g  0.41  0.45  0.89  0.15  0.32  0.47
18 h  0.88  0.68  0.69  0.96  0.23  0.27
19 i  0.94  0.57  0.64  0.90  0.14  0.86
20 j  0.05  0.10  0.99  0.69  0.41  0.05
21 k  0.53  0.90  0.66  0.80  0.41  0.44
22 l  0.89  0.25  0.71  0.02  0.37  0.80
23 m  0.55  0.04  0.54  0.48  0.15  0.12
24 > is.matrix(mat.RRR)
25 [1] TRUE
26 > is.atomic(mat.RRR)
27 [1] TRUE

```

Funkcia `diag()`

Funkcia `diag()` má veľmi zaujímavé a praktické použitie, nakoľko jej výstup závisí od toho, akú hodnotu má na vstupe. Ak je na vstupe vektor dĺžky jeden (skalárna veličina) s hodnotou n , funkcia `diag()` vytvorí jednotkovú maticu rozmeru $n \times n$:

```

1 > jednotkovaMatica.4 <- diag(4)
2 > print(jednotkovaMatica.4)
3      [,1] [,2] [,3] [,4]
4 [1,]    1    0    0    0
5 [2,]    0    1    0    0

```

```

6 [3,]    0    0    1    0
7 [4,]    0    0    0    1

```

Ak je na vstupe vektor dĺžky n , funkcia `diag()` vytvorí štvorcovú maticu rozmeru $n \times n$ s hodnotami vektora na hlavnej diagonále a všade inde budú nuly:

```

1 > print(vektor3)
2   A  B  C  D  E  F
3   3  6  9 12 15 18
4 > diag(vektor3)
5      [,1] [,2] [,3] [,4] [,5] [,6]
6 [1,]     3     0     0     0     0     0
7 [2,]     0     6     0     0     0     0
8 [3,]     0     0     9     0     0     0
9 [4,]     0     0     0    12     0     0
10 [5,]     0     0     0     0    15     0
11 [6,]     0     0     0     0     0    18

```

Ak je na vstupe štvorcová matica rozmeru $n \times n$, funkcia `diag()` vytvorí vektor dĺžky n s hodnotami hlavnej diagonály matice:

```

1 > print(postupnost1)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]     1     7    13    19    25    31
4 [2,]     2     8    14    20    26    32
5 [3,]     3     9    15    21    27    33
6 [4,]     4    10    16    22    28    34
7 [5,]     5    11    17    23    29    35
8 [6,]     6    12    18    24    30    36
9 > dim(postupnost1)
10 [1] 6 6
11 > diag(postupnost1)
12 [1]  1  8 15 22 29 36

```

Všimnime si, že funkcia funguje aj v prípade, že máme maticu rozmeru $n \times m$. Funkcia `diag()` vytvorí vektor dĺžky $\min(n,m)$ s hodnotami hlavnej diagonály matice, teda matica nemusí byť štvorcová. Uvažujme o matici A , pričom a_{ij} predstavuje hodnotu v i -tom riadku a j -tom stĺpci tejto matice. Hlavná diagonála predstavuje všetky také hodnoty a_{ij} , pre ktoré platí $i = j$. Znamená to, že na vstupe funkcie `diag()` nemusí byť iba štvorcová matica:

```

1 > print(mat.R)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 a      31    79    51    14    67    42
4 b      50    43    14    25    90    91
5 c      69    91    57    92     9    93

```

```

6 d    99    72    26     7    42     9
7 > dim(mat.R)
8 [1] 4 6
9 > diag(mat.R)
10 [1] 31 43 57 7

```

Aritmetické operácie s maticami

V lineárnej algebre často potrebujeme s maticami robiť rôzne operácie, ako napríklad transpozíciu, inverziu, výpočet determinantu alebo súčinu matic. Mnohé z týchto operácií sú ručným výpočtom mimoriadne zdĺhavé (napríklad výpočet inverznej matice), a preto je veľmi vhodné používať pri tom špecializovaný softvér. **R** v tomto smere ponúka veľmi dobré nástroje na efektívnu prácu s maticami a výrazným spôsobom tak uľahčuje výpočty.

Začnime najjednoduchšími aritmetickými operáciami. Operátory súčtu (+), rozdielu (-), násobenia (*) a delenia (/) matic predstavujú **súčet**, **rozdiel**, **súčin**, respektíve **podiel** jednotlivých odpovedajúcich si **prvkov** daných **matic**. Matice preto musia mať **rovnaký rozmer**.

```

1 > mat.A <- matrix(seq(
2 +   from = 15,
3 +   by = -1,
4 +   length.out = 24
5 + ), nrow = 4)
6 > print(mat.A)
7      [,1] [,2] [,3] [,4] [,5] [,6]
8 [1,]   15   11    7    3   -1   -5
9 [2,]   14   10    6    2   -2   -6
10 [3,]   13    9    5    1   -3   -7
11 [4,]   12    8    4    0   -4   -8
12 > dim(mat.A)
13 [1] 4 6
14 >
15 > print(mat.R)
16      [,1] [,2] [,3] [,4] [,5] [,6]
17 a    31    79    51    14    67    42
18 b    50    43    14    25    90    91
19 c    69    91    57    92     9    93
20 d    99    72    26     7    42     9
21 > dim(mat.R)
22 [1] 4 6
23 >
24 > mat.A + mat.R
25      [,1] [,2] [,3] [,4] [,5] [,6]
26 a    46    90    58    17    66    37

```



```

27 b    64    53    20    27    88    85
28 c    82   100    62    93     6    86
29 d   111    80    30     7    38     1
30 > mat.A - mat.R
31    [,1] [,2] [,3] [,4] [,5] [,6]
32 a   -16  -68  -44  -11  -68  -47
33 b   -36  -33   -8  -23  -92  -97
34 c   -56  -82  -52  -91  -12 -100
35 d   -87  -64  -22   -7  -46  -17
36 > mat.A * mat.R
37    [,1] [,2] [,3] [,4] [,5] [,6]
38 a   465   869   357    42   -67  -210
39 b   700   430    84    50  -180  -546
40 c   897   819   285    92   -27  -651
41 d  1188   576   104     0  -168   -72
42 > mat.A / mat.R
43    [,1]      [,2]      [,3]      [,4]      [,5]      [,6]
44 a 0.4838710 0.1392405 0.1372549 0.21428571 -0.01492537 -0.11904762
45 b 0.2800000 0.2325581 0.4285714 0.08000000 -0.02222222 -0.06593407
46 c 0.1884058 0.0989011 0.0877193 0.01086957 -0.33333333 -0.07526882
47 d 0.1212121 0.1111111 0.1538462 0.00000000 -0.09523810 -0.88888889

```

Všimnime si, že dané operácie naozaj nefungujú v prípade, ak matice nemajú rovnaký rozmer:

```

1 > print(mat.A)
2    [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]   15   11    7    3   -1   -5
4 [2,]   14   10    6    2   -2   -6
5 [3,]   13    9    5    1   -3   -7
6 [4,]   12    8    4    0   -4   -8
7 > dim(mat.A)
8 [1] 4 6
9 > print(postupnost1)
10    [,1] [,2] [,3] [,4] [,5] [,6]
11 [1,]    1    7   13   19   25   31
12 [2,]    2    8   14   20   26   32
13 [3,]    3    9   15   21   27   33
14 [4,]    4   10   16   22   28   34
15 [5,]    5   11   17   23   29   35
16 [6,]    6   12   18   24   30   36
17 > dim(postupnost1)
18 [1] 6 6
19 > mat.A + postupnost1
20 Error in mat.A + postupnost1 : non-conformable arrays

```

Predchádzajúce aritmetické operácie fungujú aj v kombinácii matica a skalár (vektor dĺžky 1). Vtedy je príslušná operácia vykonaná pre každý jeden prvok matice a tento skalár napríklad:

```

1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]   15   11    7    3   -1   -5
4 [2,]   14   10    6    2   -2   -6
5 [3,]   13    9    5    1   -3   -7
6 [4,]   12    8    4    0   -4   -8
7 > mat.A + 5
8      [,1] [,2] [,3] [,4] [,5] [,6]
9 [1,]   20   16   12    8    4    0
10 [2,]   19   15   11    7    3   -1
11 [3,]   18   14   10    6    2   -2
12 [4,]   17   13    9    5    1   -3
13 > mat.A - 10
14      [,1] [,2] [,3] [,4] [,5] [,6]
15 [1,]     5     1   -3   -7  -11  -15
16 [2,]     4     0   -4   -8  -12  -16
17 [3,]     3    -1   -5   -9  -13  -17
18 [4,]     2    -2   -6  -10  -14  -18
19 > mat.A * 3
20      [,1] [,2] [,3] [,4] [,5] [,6]
21 [1,]   45   33   21    9   -3  -15
22 [2,]   42   30   18    6   -6  -18
23 [3,]   39   27   15    3   -9  -21
24 [4,]   36   24   12    0  -12  -24
25 > mat.A / 2
26      [,1] [,2] [,3] [,4] [,5] [,6]
27 [1,]   7.5  5.5  3.5  1.5 -0.5 -2.5
28 [2,]   7.0  5.0  3.0  1.0 -1.0 -3.0
29 [3,]   6.5  4.5  2.5  0.5 -1.5 -3.5
30 [4,]   6.0  4.0  2.0  0.0 -2.0 -4.0

```

Ďalšou aritmetickou operáciou je umocňovanie, ktoré taktiež funguje po prvkoch:

```

1 > print(postupnost1)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]     1     7    13    19    25    31
4 [2,]     2     8    14    20    26    32
5 [3,]     3     9    15    21    27    33
6 [4,]     4    10    16    22    28    34
7 [5,]     5    11    17    23    29    35
8 [6,]     6    12    18    24    30    36
9 > postupnost1 ^ 2
10      [,1] [,2] [,3] [,4] [,5] [,6]

```

11	[1,]	1	49	169	361	625	961
12	[2,]	4	64	196	400	676	1024
13	[3,]	9	81	225	441	729	1089
14	[4,]	16	100	256	484	784	1156
15	[5,]	25	121	289	529	841	1225
16	[6,]	36	144	324	576	900	1296

Skutočné **násobenie matic** vykonávame prostredníctvom operátora `%%`, ktorý sme využívali aj pri skalárnom súčine vektorov.

Maticu A možno vynásobiť maticou B sprava vtedy, ak počet stĺpcov matice A sa rovná počtu riadkov matice B . Postup je nasledovný: Ak $A = [a_{ik}]_{m \times \ell}$ a $B = [b_{kj}]_{\ell \times n}$ sú matice rozmeru $m \times \ell$ a $\ell \times n$, ich súčinom je matica $C = AB$ rozmeru $m \times n$, ktorej hodnotu v i -tom riadku a j -tom stĺpci vypočítame podľa vzťahu:

$$c_{ij} = a_{i1}b_{1j} + a_{i2}b_{2j} + \dots + a_{i\ell}b_{\ell j} \quad (6.2)$$

Hodnota v i -tom riadku a j -tom stĺpci matice C je skalárnym súčinom i -teho riadku matice A a j -teho stĺpca matice B .

```

1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]  15   11    7    3   -1   -5
4 [2,]  14   10    6    2   -2   -6
5 [3,]  13    9    5    1   -3   -7
6 [4,]  12    8    4    0   -4   -8
7 > dim(mat.A)
8 [1] 4 6
9 > print(postupnost1)
10     [,1] [,2] [,3] [,4] [,5] [,6]
11 [1,]    1    7   13   19   25   31
12 [2,]    2    8   14   20   26   32
13 [3,]    3    9   15   21   27   33
14 [4,]    4   10   16   22   28   34
15 [5,]    5   11   17   23   29   35
16 [6,]    6   12   18   24   30   36
17 > dim(postupnost1)
18 [1] 6 6
19 > mat.A %% postupnost1
20     [,1] [,2] [,3] [,4] [,5] [,6]
21 [1,]   35  215  395  575  755  935
22 [2,]   14  158  302  446  590  734
23 [3,]   -7  101  209  317  425  533
24 [4,]  -28   44  116  188  260  332
25 > postupnost1 %% mat.A
26 Error in postupnost1 %% mat.A : non-conformable arguments

```

Vidíme, že ak násobíme maticu rozmeru 4×6 maticou rozmeru 6×6 , výsledná matica bude rozmeru 6×6 . Ak však násobíme maticu rozmeru 6×6 maticou rozmeru 4×6 , tak prvá matica nemá toľko stĺpcov, ako má druhá matica riadkov, preto násobenie týchto dvoch matíc nie je možné, na čo nás aj **R** upozorní chybovou hláškou.

Môžeme si ukázať, že ak násobíme štvorcovú maticu zľava alebo sprava jednotkovou maticou, dostaneme pôvodnú maticu:

```

1 > print(postupnost1)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]    1    7   13   19   25   31
4 [2,]    2    8   14   20   26   32
5 [3,]    3    9   15   21   27   33
6 [4,]    4   10   16   22   28   34
7 [5,]    5   11   17   23   29   35
8 [6,]    6   12   18   24   30   36
9 > dim(postupnost1)
10 [1] 6 6
11 > diag(6)
12      [,1] [,2] [,3] [,4] [,5] [,6]
13 [1,]    1    0    0    0    0    0
14 [2,]    0    1    0    0    0    0
15 [3,]    0    0    1    0    0    0
16 [4,]    0    0    0    1    0    0
17 [5,]    0    0    0    0    1    0
18 [6,]    0    0    0    0    0    1
19 > postupnost1 %*% diag(6)
20      [,1] [,2] [,3] [,4] [,5] [,6]
21 [1,]    1    7   13   19   25   31
22 [2,]    2    8   14   20   26   32
23 [3,]    3    9   15   21   27   33
24 [4,]    4   10   16   22   28   34
25 [5,]    5   11   17   23   29   35
26 [6,]    6   12   18   24   30   36
27 > diag(6) %*% postupnost1
28      [,1] [,2] [,3] [,4] [,5] [,6]
29 [1,]    1    7   13   19   25   31
30 [2,]    2    8   14   20   26   32
31 [3,]    3    9   15   21   27   33
32 [4,]    4   10   16   22   28   34
33 [5,]    5   11   17   23   29   35
34 [6,]    6   12   18   24   30   36

```

Často využívanou aritmetickou operáciou je **transpozícia matice**. Zjednodušene povedané ide o takú operáciu, pri ktorej sa riadky pôvodnej matice stávajú stĺpcami v transpo-

novanej matici a naopak. Pre transpozíciu v **R** používame funkciu `t()`, ktorej argumentom je matica, ktorú chceme transponovať. Je zrejmé, že transpozíciou transponovanej matice dostávame pôvodnú maticu:

$$(A^T)^T = A \quad (6.3)$$

```

1 > print(mat.RRR)
2      A      B      C      D      E      F
3 a 31.00 79.00 51.00 14.00 67.00 42.00
4 b 50.00 43.00 14.00 25.00 90.00 91.00
5 c 69.00 91.00 57.00 92.00  9.00 93.00
6 d 99.00 72.00 26.00  7.00 42.00  9.00
7 e  0.29  0.46  0.33  0.59  0.76  0.14
8 f  0.79  0.96  0.95  0.29  0.22  0.23
9 g  0.41  0.45  0.89  0.15  0.32  0.47
10 h  0.88  0.68  0.69  0.96  0.23  0.27
11 i  0.94  0.57  0.64  0.90  0.14  0.86
12 j  0.05  0.10  0.99  0.69  0.41  0.05
13 k  0.53  0.90  0.66  0.80  0.41  0.44
14 l  0.89  0.25  0.71  0.02  0.37  0.80
15 m  0.55  0.04  0.54  0.48  0.15  0.12
16 > t(mat.RRR)
17      a  b  c  d      e      f      g      h      i      j      k      l      m
18 A 31 50 69 99 0.29 0.79 0.41 0.88 0.94 0.05 0.53 0.89 0.55
19 B 79 43 91 72 0.46 0.96 0.45 0.68 0.57 0.10 0.90 0.25 0.04
20 C 51 14 57 26 0.33 0.95 0.89 0.69 0.64 0.99 0.66 0.71 0.54
21 D 14 25 92  7 0.59 0.29 0.15 0.96 0.90 0.69 0.80 0.02 0.48
22 E 67 90  9 42 0.76 0.22 0.32 0.23 0.14 0.41 0.41 0.37 0.15
23 F 42 91 93  9 0.14 0.23 0.47 0.27 0.86 0.05 0.44 0.80 0.12
24 > t(t(mat.RRR))
25      A      B      C      D      E      F
26 a 31.00 79.00 51.00 14.00 67.00 42.00
27 b 50.00 43.00 14.00 25.00 90.00 91.00
28 c 69.00 91.00 57.00 92.00  9.00 93.00
29 d 99.00 72.00 26.00  7.00 42.00  9.00
30 e  0.29  0.46  0.33  0.59  0.76  0.14
31 f  0.79  0.96  0.95  0.29  0.22  0.23
32 g  0.41  0.45  0.89  0.15  0.32  0.47
33 h  0.88  0.68  0.69  0.96  0.23  0.27
34 i  0.94  0.57  0.64  0.90  0.14  0.86
35 j  0.05  0.10  0.99  0.69  0.41  0.05
36 k  0.53  0.90  0.66  0.80  0.41  0.44
37 l  0.89  0.25  0.71  0.02  0.37  0.80
38 m  0.55  0.04  0.54  0.48  0.15  0.12
39 > all.equal(mat.RRR, t(t(mat.RRR)))

```

```
40 [1] TRUE
```

Na príklade výpočtov v **R** si môžeme ukázať aj ďalšie vlastnosti transpozície matice. Na overenie, či ide o rovnaké matice, môžeme používať funkciu `all.equal()`, ktorej argumenty sú matice, ktoré chceme porovnať. Pomocou tejto funkcie je možné porovnať objekty, ktoré sú úplne rovnaké alebo takmer rovnaké (berie do úvahy aj zaokrúhľovaciu chybu dátového typu `double`, ktorú sme uvádzali v Kapitole 3).

Pri transpozícii súčinu konštanty a matice môžeme konštantu vybrať pred zátvorku:

$$(kA)^T = k(A^T) \quad (6.4)$$

```
1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]    15    11     7     3    -1    -5
4 [2,]    14    10     6     2    -2    -6
5 [3,]    13     9     5     1    -3    -7
6 [4,]    12     8     4     0    -4    -8
7 > 5 * mat.A
8      [,1] [,2] [,3] [,4] [,5] [,6]
9 [1,]    75    55    35    15    -5   -25
10 [2,]    70    50    30    10   -10   -30
11 [3,]    65    45    25     5   -15   -35
12 [4,]    60    40    20     0   -20   -40
13 > t(5 * mat.A)
14      [,1] [,2] [,3] [,4]
15 [1,]    75    70    65    60
16 [2,]    55    50    45    40
17 [3,]    35    30    25    20
18 [4,]    15    10     5     0
19 [5,]    -5   -10   -15   -20
20 [6,]   -25   -30   -35   -40
21 >
22 > t(mat.A)
23      [,1] [,2] [,3] [,4]
24 [1,]    15    14    13    12
25 [2,]    11    10     9     8
26 [3,]     7     6     5     4
27 [4,]     3     2     1     0
28 [5,]    -1    -2    -3    -4
29 [6,]    -5    -6    -7    -8
30 > 5 * t(mat.A)
31      [,1] [,2] [,3] [,4]
32 [1,]    75    70    65    60
33 [2,]    55    50    45    40
```

```

34 [3,] 35 30 25 20
35 [4,] 15 10 5 0
36 [5,] -5 -10 -15 -20
37 [6,] -25 -30 -35 -40
38 >
39 > all.equal(t(5 * mat.A), 5 * t(mat.A))
40 [1] TRUE

```

Transpozícia súčtu matíc je rovná súčtu transponovaných matíc:

$$(A + B)^T = A^T + B^T \quad (6.5)$$

```

1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,] 15 11 7 3 -1 -5
4 [2,] 14 10 6 2 -2 -6
5 [3,] 13 9 5 1 -3 -7
6 [4,] 12 8 4 0 -4 -8
7 > print(mat.R)
8      [,1] [,2] [,3] [,4] [,5] [,6]
9 a 31 79 51 14 67 42
10 b 50 43 14 25 90 91
11 c 69 91 57 92 9 93
12 d 99 72 26 7 42 9
13 > mat.A + mat.R
14      [,1] [,2] [,3] [,4] [,5] [,6]
15 a 46 90 58 17 66 37
16 b 64 53 20 27 88 85
17 c 82 100 62 93 6 86
18 d 111 80 30 7 38 1
19 > t(mat.A + mat.R)
20      a b c d
21 [1,] 46 64 82 111
22 [2,] 90 53 100 80
23 [3,] 58 20 62 30
24 [4,] 17 27 93 7
25 [5,] 66 88 6 38
26 [6,] 37 85 86 1
27 > t(mat.A)
28      [,1] [,2] [,3] [,4]
29 [1,] 15 14 13 12
30 [2,] 11 10 9 8
31 [3,] 7 6 5 4
32 [4,] 3 2 1 0
33 [5,] -1 -2 -3 -4

```

```

34 [6,]    -5    -6    -7    -8
35 > t(mat.R)
36      a  b  c  d
37 [1,] 31 50 69 99
38 [2,] 79 43 91 72
39 [3,] 51 14 57 26
40 [4,] 14 25 92  7
41 [5,] 67 90  9 42
42 [6,] 42 91 93  9
43 > t(mat.A) + t(mat.R)
44      a  b  c  d
45 [1,] 46 64 82 111
46 [2,] 90 53 100 80
47 [3,] 58 20 62 30
48 [4,] 17 27 93  7
49 [5,] 66 88  6 38
50 [6,] 37 85 86  1
51 > all.equal(t(mat.A + mat.R), t(mat.A) + t(mat.R))
52 [1] TRUE

```

Transpozícia súčinnu matíc je rovná súčinnu transponovaných matíc v opačnom poradí:

$$(AB)^T = B^T A^T \quad (6.6)$$

```

1 > print(mat.RRR)
2      A      B      C      D      E      F
3 a 31.00 79.00 51.00 14.00 67.00 42.00
4 b 50.00 43.00 14.00 25.00 90.00 91.00
5 c 69.00 91.00 57.00 92.00  9.00 93.00
6 d 99.00 72.00 26.00  7.00 42.00  9.00
7 e  0.29  0.46  0.33  0.59  0.76  0.14
8 f  0.79  0.96  0.95  0.29  0.22  0.23
9 g  0.41  0.45  0.89  0.15  0.32  0.47
10 h  0.88  0.68  0.69  0.96  0.23  0.27
11 i  0.94  0.57  0.64  0.90  0.14  0.86
12 j  0.05  0.10  0.99  0.69  0.41  0.05
13 k  0.53  0.90  0.66  0.80  0.41  0.44
14 l  0.89  0.25  0.71  0.02  0.37  0.80
15 m  0.55  0.04  0.54  0.48  0.15  0.12
16 > dim(mat.RRR)
17 [1] 13  6
18 > print(postupnost1)
19      [,1] [,2] [,3] [,4] [,5] [,6]
20 [1,]    1    7   13   19   25   31
21 [2,]    2    8   14   20   26   32

```



```

22 [3,]      3      9     15     21     27     33
23 [4,]      4     10     16     22     28     34
24 [5,]      5     11     17     23     29     35
25 [6,]      6     12     18     24     30     36
26 > dim(postupnost1)
27 [1] 6 6
28 >
29 > mat.RRR %*% postupnost1
30      [,1]      [,2]      [,3]      [,4]      [,5]      [,6]
31 a  985.00 2689.00 4393.00 6097.00 7801.00 9505.00
32 b 1274.00 3152.00 5030.00 6908.00 8786.00 10664.00
33 c 1393.00 3859.00 6325.00 8791.00 11257.00 13723.00
34 d  613.00 2143.00 3673.00 5203.00 6733.00 8263.00
35 e    9.20   24.62   40.04   55.46   70.88   86.30
36 f    9.20   29.84   50.48   71.12   91.76  112.40
37 g    9.00   25.14   41.28   57.42   73.56   89.70
38 h   10.92   33.18   55.44   77.70   99.96  122.22
39 i   13.46   37.76   62.06   86.36  110.66  134.96
40 j    8.33   22.07   35.81   49.55   63.29   77.03
41 k   12.20   34.64   57.08   79.52  101.96  124.40
42 l   10.25   28.49   46.73   64.97   83.21  101.45
43 m    5.64   16.92   28.20   39.48   50.76   62.04
44 > t(mat.RRR %*% postupnost1)
45      a      b      c      d      e      f      g      h      i      j
46 [1,]  985  1274  1393  613  9.20   9.20  9.00  10.92  13.46  8.33
47 [2,] 2689  3152  3859 2143 24.62  29.84 25.14  33.18  37.76 22.07
48 [3,] 4393  5030  6325 3673 40.04  50.48 41.28  55.44  62.06 35.81
49 [4,] 6097  6908  8791 5203 55.46  71.12 57.42  77.70  86.36 49.55
50 [5,] 7801  8786 11257 6733 70.88  91.76 73.56  99.96 110.66 63.29
51 [6,] 9505 10664 13723 8263 86.30 112.40 89.70 122.22 134.96 77.03
52      k      l      m
53 [1,]  12.20  10.25  5.64
54 [2,]  34.64  28.49 16.92
55 [3,]  57.08  46.73 28.20
56 [4,]  79.52  64.97 39.48
57 [5,] 101.96  83.21 50.76
58 [6,] 124.40 101.45 62.04
59 >
60 > t(postupnost1)
61      [,1] [,2] [,3] [,4] [,5] [,6]
62 [1,]    1    2    3    4    5    6
63 [2,]    7    8    9   10   11   12
64 [3,]   13   14   15   16   17   18
65 [4,]   19   20   21   22   23   24
66 [5,]   25   26   27   28   29   30
67 [6,]   31   32   33   34   35   36

```

```

68 > t(mat.RRR)
69   a  b  c  d  e  f  g  h  i  j  k  l  m
70 A 31 50 69 99 0.29 0.79 0.41 0.88 0.94 0.05 0.53 0.89 0.55
71 B 79 43 91 72 0.46 0.96 0.45 0.68 0.57 0.10 0.90 0.25 0.04
72 C 51 14 57 26 0.33 0.95 0.89 0.69 0.64 0.99 0.66 0.71 0.54
73 D 14 25 92 7 0.59 0.29 0.15 0.96 0.90 0.69 0.80 0.02 0.48
74 E 67 90 9 42 0.76 0.22 0.32 0.23 0.14 0.41 0.41 0.37 0.15
75 F 42 91 93 9 0.14 0.23 0.47 0.27 0.86 0.05 0.44 0.80 0.12
76 > t(postupnost1) %*% t(mat.RRR)
77   a  b  c  d  e  f  g  h  i  j
78 [1,] 985 1274 1393 613 9.20 9.20 9.00 10.92 13.46 8.33
79 [2,] 2689 3152 3859 2143 24.62 29.84 25.14 33.18 37.76 22.07
80 [3,] 4393 5030 6325 3673 40.04 50.48 41.28 55.44 62.06 35.81
81 [4,] 6097 6908 8791 5203 55.46 71.12 57.42 77.70 86.36 49.55
82 [5,] 7801 8786 11257 6733 70.88 91.76 73.56 99.96 110.66 63.29
83 [6,] 9505 10664 13723 8263 86.30 112.40 89.70 122.22 134.96 77.03
84   k  l  m
85 [1,] 12.20 10.25 5.64
86 [2,] 34.64 28.49 16.92
87 [3,] 57.08 46.73 28.20
88 [4,] 79.52 64.97 39.48
89 [5,] 101.96 83.21 50.76
90 [6,] 124.40 101.45 62.04
91 >
92 > all.equal(t(mat.RRR %*% postupnost1), t(postupnost1) %*% t(mat.RRR
93 [1] TRUE

```

Inverzná matica (označujeme A^{-1}) k matici A predstavuje maticu, ktorú keď vynásobíme pôvodnou maticou, dostaneme jednotkovú maticu. Vzhľadom na definíciu násobenia matic platí, že ako pôvodná matica, tak aj matica k nej inverzná musia byť štvorcové matice. V prostredí **R** na výpočet inverznej matice používame funkciu `solve()`.

$$A_{n \times n} A_{n \times n}^{-1} = I_n \quad (6.7)$$

```

1 > set.seed(123)
2 > mat.B <- matrix(runif(25), nrow = 5)
3 > print(mat.B)
4   [,1] [,2] [,3] [,4] [,5]
5 [1,] 0.2875775 0.0455565 0.9568333 0.89982497 0.8895393
6 [2,] 0.7883051 0.5281055 0.4533342 0.24608773 0.6928034
7 [3,] 0.4089769 0.8924190 0.6775706 0.04205953 0.6405068
8 [4,] 0.8830174 0.5514350 0.5726334 0.32792072 0.9942698
9 [5,] 0.9404673 0.4566147 0.1029247 0.95450365 0.6557058
10 > solve(mat.B)

```

```

11      [,1]      [,2]      [,3]      [,4]      [,5]
12 [1,]  0.008815086  5.4700056 -1.3791158 -2.6317602 -0.4536595
13 [2,] -0.490944922 -1.5615139  1.7258450 -0.1358457  0.8360285
14 [3,]  1.118627201  5.2806716 -0.2895156 -3.7610801 -1.1111159
15 [4,]  0.473805674 -0.4918888  0.3439892 -0.9991009  1.0559032
16 [5,] -0.536065426 -6.8709975  0.3209160  5.9140349  0.2309065
17 > mat.B %*% solve(mat.B)
18      [,1]      [,2]      [,3]      [,4]
19 [1,]  1.000000e+00  8.881784e-16  0.000000e+00 -8.881784e-16
20 [2,]  1.110223e-16  1.000000e+00 -2.220446e-16 -8.881784e-16
21 [3,]  0.000000e+00  8.881784e-16  1.000000e+00  4.440892e-16
22 [4,]  0.000000e+00  1.776357e-15 -3.330669e-16  1.000000e+00
23 [5,]  0.000000e+00  8.881784e-16 -1.387779e-16  0.000000e+00
24      [,5]
25 [1,]  2.498002e-16
26 [2,] -1.110223e-16
27 [3,] -2.775558e-17
28 [4,] -3.053113e-16
29 [5,]  1.000000e+00
30 > round(mat.B %*% solve(mat.B))
31      [,1] [,2] [,3] [,4] [,5]
32 [1,]    1    0    0    0    0
33 [2,]    0    1    0    0    0
34 [3,]    0    0    1    0    0
35 [4,]    0    0    0    1    0
36 [5,]    0    0    0    0    1
37 > diag(5)
38      [,1] [,2] [,3] [,4] [,5]
39 [1,]    1    0    0    0    0
40 [2,]    0    1    0    0    0
41 [3,]    0    0    1    0    0
42 [4,]    0    0    0    1    0
43 [5,]    0    0    0    0    1
44 > all.equal(mat.B %*% solve(mat.B), diag(5))
45 [1] TRUE

```

Aj pre inverziu matíc platia nejaké pravidlá podobne ako pre transpozíciu. Môžeme si na praktickom príklade ukázať jedno z nich a to, že inverzia súčinu matíc je rovná súčinu inverzií matíc v opačnom poradí:

$$(AB)^{-1} = B^{-1}A^{-1} \quad (6.8)$$

```

1 > set.seed(123)
2 > mat.D <- matrix(sample(1:100, 25), nrow = 5)
3 > print(mat.B)

```

```

4          [,1]      [,2]      [,3]      [,4]      [,5]
5 [1,] 0.2875775 0.0455565 0.9568333 0.89982497 0.8895393
6 [2,] 0.7883051 0.5281055 0.4533342 0.24608773 0.6928034
7 [3,] 0.4089769 0.8924190 0.6775706 0.04205953 0.6405068
8 [4,] 0.8830174 0.5514350 0.5726334 0.32792072 0.9942698
9 [5,] 0.9404673 0.4566147 0.1029247 0.95450365 0.6557058
10 > print(mat.D)
11      [,1] [,2] [,3] [,4] [,5]
12 [1,]    31   42   90   26   78
13 [2,]    79   50   69    7   93
14 [3,]    51   43   57   95   76
15 [4,]    14   97    9   87   15
16 [5,]    67   25   72   36   32
17 >
18 > mat.B %*% mat.D
19      [,1]      [,2]      [,3]      [,4]      [,5]
20 [1,] 133.5091 165.0214 155.7101 209.0033 141.3498
21 [2,] 139.1409 120.1981 185.3234 113.6100 170.9160
22 [3,] 161.2383 111.0260 183.5014 107.9670 187.5177
23 [4,] 171.3482 145.9468 224.6994 145.5415 200.4144
24 [5,] 127.7715 175.7356 177.8165 144.0735 158.9440
25 > solve(mat.B %*% mat.D)
26      [,1]      [,2]      [,3]      [,4]      [,5]
27 [1,] -0.010588921 -0.17588353 0.042319298 0.10481591 0.016457523
28 [2,] -0.008978980 -0.04677675 0.011557043 0.01355903 0.027553722
29 [3,] -0.008340289 -0.01193476 -0.033627983 0.05270152 -0.006527852
30 [4,] 0.015687550 0.04793174 -0.012658989 -0.02565164 -0.018213913
31 [5,] 0.013550494 0.16301182 0.002297979 -0.13495803 -0.013590136
32 >
33 > solve(mat.D)
34      [,1]      [,2]      [,3]      [,4]
35 [1,] -0.016355380 0.009713003 0.0005234426 -0.0009509526
36 [2,] 0.003156288 0.005259973 -0.0113658639 0.0116418023
37 [3,] 0.014837347 -0.010612269 -0.0066692618 -0.0008051140
38 [4,] -0.002661443 -0.007731391 0.0115984603 -0.0004179186
39 [5,] 0.001388320 0.008129467 0.0097411943 -0.0048224362
40      [,5]
41 [1,] 0.010840408
42 [2,] -0.001443418
43 [3,] 0.010892768
44 [4,] 0.001606179
45 [5,] -0.016635114
46 > solve(mat.B)
47      [,1]      [,2]      [,3]      [,4]      [,5]
48 [1,] 0.008815086 5.4700056 -1.3791158 -2.6317602 -0.4536595
49 [2,] -0.490944922 -1.5615139 1.7258450 -0.1358457 0.8360285

```

```

50 [3,] 1.118627201 5.2806716 -0.2895156 -3.7610801 -1.1111159
51 [4,] 0.473805674 -0.4918888 0.3439892 -0.9991009 1.0559032
52 [5,] -0.536065426 -6.8709975 0.3209160 5.9140349 0.2309065
53 > solve(mat.D) %*% solve(mat.B)
54      [,1]      [,2]      [,3]      [,4]      [,5]
55 [1,] -0.010588921 -0.17588353 0.042319298 0.10481591 0.016457523
56 [2,] -0.008978980 -0.04677675 0.011557043 0.01355903 0.027553722
57 [3,] -0.008340289 -0.01193476 -0.033627983 0.05270152 -0.006527852
58 [4,] 0.015687550 0.04793174 -0.012658989 -0.02565164 -0.018213913
59 [5,] 0.013550494 0.16301182 0.002297979 -0.13495803 -0.013590136
60 >
61 > all.equal(solve(mat.B %*% mat.D), solve(mat.D) %*% solve(mat.B))
62 [1] TRUE

```

Pre výpočet **determinantu štvorcovej matice** využívame funkciu `det()`. Môžeme vidieť, že štvorcová matica uložená v objekte `mat.D` je regulárna, a preto má nenulový determinant. Štvorcová matica uložená v objekte `postupnost1` je singulárna, a preto je jej determinant rovný nule.

```

1 > print(mat.D)
2      [,1] [,2] [,3] [,4] [,5]
3 [1,]    31    42    90    26    78
4 [2,]    79    50    69     7    93
5 [3,]    51    43    57    95    76
6 [4,]    14    97     9    87    15
7 [5,]    67    25    72    36    32
8 > dim(mat.D)
9 [1] 5 5
10 > det(mat.D)
11 [1] 2127866400
12 >
13 > print(postupnost1)
14      [,1] [,2] [,3] [,4] [,5] [,6]
15 [1,]     1     7    13    19    25    31
16 [2,]     2     8    14    20    26    32
17 [3,]     3     9    15    21    27    33
18 [4,]     4    10    16    22    28    34
19 [5,]     5    11    17    23    29    35
20 [6,]     6    12    18    24    30    36
21 > dim(postupnost1)
22 [1] 6 6
23 > det(postupnost1)
24 [1] 0

```

Pre **výpočet vlastných čísel a vlastných vektorov matice** využívame funkciu `eigen()`, ktorej argumentom je štvorcová matica. Výstupom funkcie je objekt typu list, na ktorého

jednotlivé prvky sa môžeme odvolávať prostredníctvom symbolu \$:

```
1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]   15   11    7    3   -1   -5
4 [2,]   14   10    6    2   -2   -6
5 [3,]   13    9    5    1   -3   -7
6 [4,]   12    8    4    0   -4   -8
7 > dim(mat.A)
8 [1] 4 6
9 > eigen(mat.A)
10 Error in eigen(mat.A) : non-square matrix in 'eigen'
11 >
12 > print(mat.D)
13      [,1] [,2] [,3] [,4] [,5]
14 [1,]   31   42   90   26   78
15 [2,]   79   50   69    7   93
16 [3,]   51   43   57   95   76
17 [4,]   14   97    9   87   15
18 [5,]   67   25   72   36   32
19 > dim(mat.D)
20 [1] 5 5
21 > eigen(mat.D)
22 eigen() decomposition
23 $values
24 [1] 268.75410+ 0.00000i  31.23791+66.17556i  31.23791-66.17556i
25 [4] -37.11496+10.05014i -37.11496-10.05014i
26
27 $vectors
28      [,1]      [,2]      [,3]
29 [1,] -0.4503157+0i -0.2636881+0.0192168i -0.2636881-0.0192168i
30 [2,] -0.5020081+0i -0.3162006+0.4844264i -0.3162006-0.4844264i
31 [3,] -0.5123621+0i  0.1036109-0.3531097i  0.1036109+0.3531097i
32 [4,] -0.3602453+0i  0.6486444+0.0000000i  0.6486444+0.0000000i
33 [5,] -0.3910397+0i -0.1826119-0.0770663i -0.1826119+0.0770663i
34      [,4]      [,5]
35 [1,] -0.2738056+0.4387805i -0.2738056-0.4387805i
36 [2,] -0.1482730-0.1751180i -0.1482730+0.1751180i
37 [3,] -0.3637546-0.3147498i -0.3637546+0.3147498i
38 [4,]  0.0846355+0.1170436i  0.0846355-0.1170436i
39 [5,]  0.6539142+0.0000000i  0.6539142+0.0000000i
40
41 > eigen(mat.D)$values
42 [1] 268.75410+ 0.00000i  31.23791+66.17556i  31.23791-66.17556i
43 [4] -37.11496+10.05014i -37.11496-10.05014i
44 > eigen(mat.D)$vectors
45      [,1]      [,2]      [,3]
```

```

46 [1,] -0.4503157+0i -0.2636881+0.0192168i -0.2636881-0.0192168i
47 [2,] -0.5020081+0i -0.3162006+0.4844264i -0.3162006-0.4844264i
48 [3,] -0.5123621+0i 0.1036109-0.3531097i 0.1036109+0.3531097i
49 [4,] -0.3602453+0i 0.6486444+0.0000000i 0.6486444+0.0000000i
50 [5,] -0.3910397+0i -0.1826119-0.0770663i -0.1826119+0.0770663i
51                                     [,4]                                     [,5]
52 [1,] -0.2738056+0.4387805i -0.2738056-0.4387805i
53 [2,] -0.1482730-0.1751180i -0.1482730+0.1751180i
54 [3,] -0.3637546-0.3147498i -0.3637546+0.3147498i
55 [4,] 0.0846355+0.1170436i 0.0846355-0.1170436i
56 [5,] 0.6539142+0.0000000i 0.6539142+0.0000000i

```

Riadkové a stĺpcové funkcie

Na hodnoty matice môžeme použiť rôzne funkcie, napríklad ak chceme vypočítať priemernú hodnotu v matici, využívame funkciu `mean()`, ktorej argumentom je samotná matica. Existujú však aj funkcie, ktoré nám umožňujú vypočítať niektoré charakteristiky, ako napríklad súčet alebo priemer pre každý riadok alebo stĺpec matice osobitne. Ide o funkcie:

- `rowSums()`: spočíta všetky hodnoty za každý riadok matice,
- `colSums()`: spočíta všetky hodnoty za každý stĺpec matice,
- `rowMeans()`: vypočíta priemer hodnôt za každý riadok matice,
- `colMeans()`: vypočíta priemer hodnôt za každý stĺpec matice.

```

1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]  15  11   7   3  -1  -5
4 [2,]  14  10   6   2  -2  -6
5 [3,]  13   9   5   1  -3  -7
6 [4,]  12   8   4   0  -4  -8
7 > mat.A[1, ]
8 [1] 15 11  7  3 -1 -5
9 > sum(mat.A[1, ])
10 [1] 30
11 > mean(mat.A[1, ])
12 [1] 5
13 > mat.A[2, ]
14 [1] 14 10  6  2 -2 -6
15 > sum(mat.A[2, ])
16 [1] 24
17 > mean(mat.A[2, ])

```

```

18 [1] 4
19 > mat.A[3, ]
20 [1] 13 9 5 1 -3 -7
21 > sum(mat.A[3, ])
22 [1] 18
23 > mean(mat.A[3, ])
24 [1] 3
25 > mat.A[4, ]
26 [1] 12 8 4 0 -4 -8
27 > sum(mat.A[4, ])
28 [1] 12
29 > mean(mat.A[4, ])
30 [1] 2
31 > rowSums(mat.A)
32 [1] 30 24 18 12
33 > rowMeans(mat.A)
34 [1] 5 4 3 2

```

```

1 > print(mat.A)
2      [,1] [,2] [,3] [,4] [,5] [,6]
3 [1,]  15   11    7    3   -1   -5
4 [2,]  14   10    6    2   -2   -6
5 [3,]  13    9    5    1   -3   -7
6 [4,]  12    8    4    0   -4   -8
7 > colSums(mat.A)
8 [1]  54  38  22    6 -10 -26
9 > colMeans(mat.A)
10 [1] 13.5  9.5  5.5  1.5 -2.5 -6.5

```

Ak by sme chceli vypočítať napríklad riadkové mediány alebo stĺpcové rozptyly, v **R** neexistuje funkcia `rowMedians()`, alebo `colVars()`. V takýchto prípadoch je veľmi vhodné využiť funkciu `apply()`, ktorá je podrobne vysvetlená v pokračovaní tejto knihy (***R** snadno a rychle 2 Vizualizace dat a programování*).

6.1.3 Atomické pole

Maticu je možné zovšeobecniť do viacerých rozmerov. V prostredí **R** túto štruktúru nazývame pole (array). Je to vlastne viacrozmerná tabuľka, ktorej všetky prvky majú rovnaký dátový typ. Atomické polia sa vytvárajú pomocou funkcie `array()`. Pomocou tejto funkcie je možné vytvoriť aj atomickú maticu:

```

1 > mat.E <- array(1:12, dim = c(3, 4))
2 > print(mat.E)
3      [,1] [,2] [,3] [,4]
4 [1,]    1    4    7   10

```



```

5 [2,]    2    5    8   11
6 [3,]    3    6    9   12
7 > dim(mat.E)
8 [1] 3 4
9 > is.matrix(mat.E)
10 [1] TRUE
11 > is.array(mat.E)
12 [1] TRUE
13 > is.atomic(mat.E)
14 [1] TRUE

```

Pre vytvorenie poľa musíme nadefinovať aj tretiu dimenziu:

```

1 > pole <- array(1:24, dim = c(2, 3, 4))
2 > print(pole)
3 , , 1
4
5      [,1] [,2] [,3]
6 [1,]    1    3    5
7 [2,]    2    4    6
8
9 , , 2
10
11     [,1] [,2] [,3]
12 [1,]    7    9   11
13 [2,]    8   10   12
14
15 , , 3
16
17     [,1] [,2] [,3]
18 [1,]   13   15   17
19 [2,]   14   16   18
20
21 , , 4
22
23     [,1] [,2] [,3]
24 [1,]   19   21   23
25 [2,]   20   22   24
26
27 > dim(pole)
28 [1] 2 3 4
29 > is.matrix(pole)
30 [1] FALSE
31 > is.array(pole)
32 [1] TRUE
33 > is.atomic(pole)

```

```
34 [1] TRUE
```

Odkazovanie prostredníctvom indexov je rovnaké ako v prípade matice alebo vektora, je však potrebné uvádzať poradie v troch dimenziách a tieto dimenzie si nepomýliť, lebo sa ľahko môže stať, že sa dostaneme mimo poľa:

```
1 > print(pole)
2 , , 1
3
4      [,1] [,2] [,3]
5 [1,]    1    3    5
6 [2,]    2    4    6
7
8 , , 2
9
10     [,1] [,2] [,3]
11 [1,]    7    9   11
12 [2,]    8   10   12
13
14 , , 3
15
16     [,1] [,2] [,3]
17 [1,]   13   15   17
18 [2,]   14   16   18
19
20 , , 4
21
22     [,1] [,2] [,3]
23 [1,]   19   21   23
24 [2,]   20   22   24
25
26 > dim(pole)
27 [1] 2 3 4
28 > pole[1, 2, 3]
29 [1] 15
30 > pole[1, ,]
31      [,1] [,2] [,3] [,4]
32 [1,]    1    7   13   19
33 [2,]    3    9   15   21
34 [3,]    5   11   17   23
35 > pole[, 2, ]
36     [,1] [,2] [,3] [,4]
37 [1,]    3    9   15   21
38 [2,]    4   10   16   22
39 > pole[, , 4]
40     [,1] [,2] [,3]
```

```

41 [1,] 19 21 23
42 [2,] 20 22 24
43 > pole[3,1,4]
44 Error in pole[3, 1, 4] : subscript out of bounds

```

Redukovaním dimenzií vieme veľmi jednoducho z poľa urobiť maticu alebo z matice vektor:

```

1 > print(pole)
2 , , 1
3
4      [,1] [,2] [,3]
5 [1,]    1    3    5
6 [2,]    2    4    6
7
8 , , 2
9
10     [,1] [,2] [,3]
11 [1,]    7    9   11
12 [2,]    8   10   12
13
14 , , 3
15
16     [,1] [,2] [,3]
17 [1,]   13   15   17
18 [2,]   14   16   18
19
20 , , 4
21
22     [,1] [,2] [,3]
23 [1,]   19   21   23
24 [2,]   20   22   24
25
26 >
27 > dim(pole) <- c(4, 6) #zmena z troch dimenzii (pole) na dve
    dimenzie (matica)
28 > print(pole)
29      [,1] [,2] [,3] [,4] [,5] [,6]
30 [1,]    1    5    9   13   17   21
31 [2,]    2    6   10   14   18   22
32 [3,]    3    7   11   15   19   23
33 [4,]    4    8   12   16   20   24
34 >
35 > dim(pole) <- NULL #zmena z dvoch dimenzii na vektor
36 > print(pole)
37 [1] 1 2 3 4 5 6 7 8 9 10 11 12 13 14 15 16 17 18 19 20

```

6.2 Nehomogénne dátové štruktúry

6.2.1 Zoznam (list)

Zoznam (list) predstavuje nehomogénnu dátovú štruktúru a ide teda o neatomický vektor. Prvky listu môžu mať rôzne dátové typy, na prvom mieste listu môže byť uložená napríklad matica, na druhom vektor, na treťom nejaký iný list a podobne. Listy v prostredí **R** vytvárame pomocou funkcie `list()`.

Použitie listu si ukážeme na praktickom príklade. Máme troch respondentov Janu, Jakuba a Terezu, pričom sme sa ich opýtali na ich vek, hmotnosť a či majú nejakých súrodencov. List vytvorený z ich odpovedí uložený do objektu `novyList` by mohol vyzeráť nasledovne:

```

1 > novyList = list(
2 +   "meno" = c("Jana", "Jakub", "Tereza"),
3 +   "vek" = c(28L, 42L, 19L),
4 +   "hmotnost" = c(66.7, 91.5, 56.8),
5 +   "surodenci" = c(TRUE, FALSE, TRUE)
6 + )
7 > print(novyList)
8 $meno
9 [1] "Jana" "Jakub" "Tereza"
10
11 $vek
12 [1] 28 42 19
13
14 $hmotnost
15 [1] 66.7 91.5 56.8
16
17 $surodenci
18 [1] TRUE FALSE TRUE

```

Vidíme napríklad, že Jakub má 42 rokov, nemá súrodencov a váži 91,5 kg. Na prvý pohľad je jasné, že mená respondentov sú iného dátového typu ako napríklad ich hmotnosti. Ak chceme poznať dátové typy jednotlivých prvkov listu, môžeme na to použiť funkciu `str()`, ktorá slúži na zobrazenie štruktúry objektu uloženého v **R**:

```

1 > str(novyList)
2 List of 4
3  $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4  $ vek       : int [1:3] 28 42 19
5  $ hmotnost  : num [1:3] 66.7 91.5 56.8

```

```
6 $ surodenci: logi [1:3] TRUE FALSE TRUE
```

V liste sme vytvorili štyri prvky, pričom každý z týchto prvkov je vektor. Každý z nich je vektorom inej triedy:

- **meno** je vektor triedy `character`,
- **vek** je vektor triedy `integer`,
- **hmotnost** je vektor triedy `numeric`,
- **surodenci** je vektor triedy `logical`.

Na nehomogénne dátové štruktúry sa môžeme odkazovať aj prostredníctvom symbolu `$`, za ktorým nasleduje názov prvku tejto štruktúry. Pokiaľ každý z prvkov začína iným písmenom (ako je tomu v našom prípade: `meno`, `vek`, `hmotnost`, `surodenci`), potom postačí za symbolom `$` napísať začiatkové písmeno názvu prvku:

```
1 > novyList$meno
2 [1] "Jana"      "Jakub"      "Tereza"
3 > novyList$m
4 [1] "Jana"      "Jakub"      "Tereza"
```

Odkazovanie je, samozrejme, možné aj prostredníctvom poradí, pri listoch však využijeme dvojité hranatú zátvorku (`[[]]`), ak chceme na výstupe získať vektor, ktorý je v príslušnom prvku uložený. Chceme napríklad pracovať s vektorom hmotností, ktorý je tretí v poradí. Všimnime si, že notácia s jednou hranatou zátvorkou na výstupe nedáva vektor, ale list, preto pri práci s listami odporúčame nezabudnúť na používanie dvojitých hranatých zátvoriek.

```
1 > novyList$hmotnost #vektor
2 [1] 66.7 91.5 56.8
3 > novyList[[3]] #vektor
4 [1] 66.7 91.5 56.8
5 > novyList[3] #list
6 $hmotnost
7 [1] 66.7 91.5 56.8
```

Samozrejme, môžeme pracovať s viacerými prvkami súčasne, je však potrebné ich vyberať pomocou funkcie `c()`, čo sme využívali už pri práci s atomickými vektormi a maticami. Povedzme, že chceme poznať iba meno a hmotnosť respondentov:

```
1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ vek       : int [1:3] 28 42 19
```

```

5 $ hmotnost : num [1:3] 66.7 91.5 56.8
6 $ surodenci: logi [1:3] TRUE FALSE TRUE
7 > novyList[c(1, 3)]
8 $meno
9 [1] "Jana" "Jakub" "Tereza"
10
11 $hmotnost
12 [1] 66.7 91.5 56.8
13
14 > novyList[c("meno", "hmotnost")]
15 $meno
16 [1] "Jana" "Jakub" "Tereza"
17
18 $hmotnost
19 [1] 66.7 91.5 56.8

```

Vzhľadom na to, že aj samotné prvky listu, keďže ide o vektory, majú nejaké prvky, je možné sa odkazovať napríklad na tretiu hodnotu vektora `meno` nasledujúcim spôsobom:

```

1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ vek       : int [1:3] 28 42 19
5 $ hmotnost  : num [1:3] 66.7 91.5 56.8
6 $ surodenci: logi [1:3] TRUE FALSE TRUE
7 > novyList$meno[3]
8 [1] "Tereza"
9 > novyList[[1]][3]
10 [1] "Tereza"
11 > novyList[[c(1, 3)]]
12 [1] "Tereza"

```

Ak by sme chceli napríklad zistiť, koľko rokov má Jakub, tak môžeme využiť kombináciu vyššie uvedených funkcií a relačného operátora `==` (je rovné?):

```

1 > novyList$vek[novyList$meno == "Jakub"]
2 [1] 42

```

Pre odstránenie niektorého z prvkov listu môžeme použiť priradenie hodnoty `NULL` alebo hranaté zátvorky s mínusovým znamienkom. To znamená, že daný prvok listu nechceme zobrazovať. Ak napríklad nechceme zobraziť vek, ktorý je druhý prvok listu, môžeme to urobiť nasledovne:

```

1 > novyList[-2]
2 $meno
3 [1] "Jana" "Jakub" "Tereza"

```

```

4
5 $hmotnost
6 [1] 66.7 91.5 56.8
7
8 $surodenci
9 [1] TRUE FALSE TRUE
10
11 > str(novyList)
12 List of 4
13 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
14 $ vek       : int [1:3] 28 42 19
15 $ hmotnost  : num [1:3] 66.7 91.5 56.8
16 $ surodenci: logi [1:3] TRUE FALSE TRUE

```

Vidíme, že príkaz [-2] druhý prvok iba nevybral, ale nevymazal ho, stále tam ten prvok `vek` máme. Ak by sme ho chceli vymazať, môžeme použiť priradenie hodnoty `NULL` do tohto prvku:

```

1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ vek       : int [1:3] 28 42 19
5 $ hmotnost  : num [1:3] 66.7 91.5 56.8
6 $ surodenci: logi [1:3] TRUE FALSE TRUE
7 > novyList$vek<-NULL
8 > str(novyList)
9 List of 3
10 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
11 $ hmotnost  : num [1:3] 66.7 91.5 56.8
12 $ surodenci: logi [1:3] TRUE FALSE TRUE

```

Pridanie prvku do listu je veľmi jednoduché, môžeme opäť použiť symbol `$` a za ním názov prvku, za ktorým nasleduje definícia hodnôt tohto prvku. Pridajme logický vektor s názvom `muz`, ktorý nadobúda hodnotu `TRUE` v prípade, ak je respondent muž:

```

1 > str(novyList)
2 List of 3
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ hmotnost  : num [1:3] 66.7 91.5 56.8
5 $ surodenci: logi [1:3] TRUE FALSE TRUE
6 > novyList$muz <- c(FALSE, TRUE, FALSE)
7 > str(novyList)
8 List of 4
9 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
10 $ hmotnost  : num [1:3] 66.7 91.5 56.8
11 $ surodenci: logi [1:3] TRUE FALSE TRUE

```

```
12 $ muz      : logi [1:3] FALSE TRUE FALSE
```

V tejto časti upozorňujeme na situáciu, keď máme v liste dva prvky začínajúce na písmeno m (meno a muz). Ak by sme chceli pracovať iba s jedným prvkom, v tomto prípade by nám notácia \$ v kombinácii so začiatočným písmenom vypisovala chybu, lebo prvok nie je zadaný jednoznačne. V tomto prípade postačí zobrať prvé dve začiatočné písmená:

```
1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ hmotnost  : num [1:3] 66.7 91.5 56.8
5 $ surodenci: logi [1:3] TRUE FALSE TRUE
6 $ muz       : logi [1:3] FALSE TRUE FALSE
7 > novyList$m
8 NULL
9 > novyList$me
10 [1] "Jana" "Jakub" "Tereza"
11 > novyList$meno
12 [1] "Jana" "Jakub" "Tereza"
13 > novyList$mu
14 [1] FALSE TRUE FALSE
15 > novyList$muz
16 [1] FALSE TRUE FALSE
```

List je veľmi flexibilná dátová štruktúra, ktorého prvky môžu byť ľubovoľné objekty, dokonca aj funkcie. Vytvorme si list funkcií pozostávajúci z funkcií, ktoré si sami vyberieme a s ktorými chceme ďalej pracovať. Vyberme si napríklad funkcie priemer, rozptyl a súčin. Objekt, ktorý vytvárame, nazvime listFunkcii:

```
1 > listFunkcii <- list("priemer" = mean,
2 +                    "rozptyl" = var,
3 +                    "sucin" = prod)
4 > str(listFunkcii)
5 List of 3
6 $ priemer:function (x, ...)
7 $ rozptyl:function (x, y = NULL, na.rm = FALSE, use)
8 $ sucin :function (... , na.rm = FALSE)
9 > print(vektor1)
10 A B C D E F G H I J
11 1 2 3 4 5 6 7 8 9 10
12 > listFunkcii$priemer(vektor1)
13 [1] 5.5
14 > mean(vektor1)
15 [1] 5.5
16 > listFunkcii$rozptyl(vektor1)
17 [1] 9.166667
```



```

18 > var(vektor1)
19 [1] 9.166667
20 > listFunkcii$sucin(vektor1)
21 [1] 3628800
22 > prod(vektor1)
23 [1] 3628800

```

Funkciu `length()` už poznáme, pri atomických objektoch nám udávala počet hodnôt (vektora, matice alebo poľa). V prípade listu táto funkcia vypočíta počet prvkov listu:

```

1 > print(vektor1)
2  A B C D E F G H I J
3  1 2 3 4 5 6 7 8 9 10
4 > length(vektor1)
5 [1] 10
6 >
7 > print(jednotkovaMatica.4)
8      [,1] [,2] [,3] [,4]
9 [1,]    1    0    0    0
10 [2,]    0    1    0    0
11 [3,]    0    0    1    0
12 [4,]    0    0    0    1
13 > length(jednotkovaMatica.4)
14 [1] 16
15 >
16 > str(novyList)
17 List of 4
18 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
19 $ hmotnost  : num [1:3] 66.7 91.5 56.8
20 $ surodenci: logi [1:3] TRUE FALSE TRUE
21 $ muz       : logi [1:3] FALSE TRUE FALSE
22 > length(novyList)
23 [1] 4
24 >
25 > str(listFunkcii)
26 List of 3
27 $ priemer:function (x, ...)
28 $ rozptyl:function (x, y = NULL, na.rm = FALSE, use)
29 $ sucin :function (... , na.rm = FALSE)
30 > length(listFunkcii)
31 [1] 3

```

Funkcia `names()`, v prípade, ak pracuje s listami, na výstupe vypíše mená prvkov listu:

```

1 > names(novyList)
2 [1] "meno"      "hmotnost"  "surodenci" "muz"
3 > names(listFunkcii)

```

```
4 [1] "priemer" "rozptyl" "sucin"
```

List ako neatomický vektor môžeme veľmi jednoducho konvertovať na atomický vektor pomocou funkcie `unlist()`, ktorá vytvorí zo všetkých hodnôt všetkých prvkov jeden dlhý atomický vektor. Samozrejme, platí pravidlo konverzie na nadradený dátový typ, ako sme to vysvetľovali pri vektoroch (viď Kapitola 6.1.1). V objekte `novyList` sa nachádzajú štyri prvky, pričom nadradenou triedou je trieda `character`. Preto aj vektor, ktorý vznikne funkciou `unlist()` z tohto listu, bude mať triedu `character`:

```
1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ hmotnost  : num [1:3] 66.7 91.5 56.8
5 $ surodenci: logi [1:3] TRUE FALSE TRUE
6 $ muz       : logi [1:3] FALSE TRUE FALSE
7 > typeof(novyList)
8 [1] "list"
9 > unlist(novyList)
10      meno1      meno2      meno3  hmotnost1  hmotnost2  hmotnost3
11      "Jana"     "Jakub"     "Tereza"    "66.7"     "91.5"     "56.8"
12 surodenci1 surodenci2 surodenci3      muz1      muz2      muz3
13      "TRUE"    "FALSE"    "TRUE"    "FALSE"    "TRUE"    "FALSE"
14 > typeof(unlist(novyList))
15 [1] "character"
```

Všimnime si, že funkcia `unlist()` automaticky vytvára vektor aj s menovkami, ktoré pomenúvava prostredníctvom názvov prvkov a čísluje prostredníctvom poradia hodnoty v jednotlivých prvkoch. Ak nemáme záujem o tieto menovky, musíme nastaviť argument `use.names = FALSE`. Vymažme si z objektu `novyList` prvok `meno`, aby sme z listu odstránili typ `character` a následne aplikujme funkciu `unlist()`. Vidíme, že teraz je novovytvorený vektor typu `double`:

```
1 > str(novyList)
2 List of 4
3 $ meno      : chr [1:3] "Jana" "Jakub" "Tereza"
4 $ hmotnost  : num [1:3] 66.7 91.5 56.8
5 $ surodenci: logi [1:3] TRUE FALSE TRUE
6 $ muz       : logi [1:3] FALSE TRUE FALSE
7 > typeof(novyList)
8 [1] "list"
9 > novyList <- novyList[-1]
10 > str(novyList)
11 List of 3
12 $ hmotnost : num [1:3] 66.7 91.5 56.8
13 $ surodenci: logi [1:3] TRUE FALSE TRUE
14 $ muz      : logi [1:3] FALSE TRUE FALSE
```

```

15 > unlist(novyList)
16   hmotnost1  hmotnost2  hmotnost3  surodenci1  surodenci2  surodenci3
17       66.7       91.5       56.8         1.0         0.0         1.0
18       muz1       muz2       muz3
19       0.0       1.0       0.0
20 > typeof(unlist(novyList))
21 [1] "double"

```

6.2.2 Data frame

V úvode Kapitoly 6.1.2 sme uviedli základné nedostatky použitia atomických matíc pri praktických analýzach. Práve z dôvodu atomickej štruktúry matice nie sú vhodné na reprezentáciu dát, ktoré chceme analyzovať, nakoľko tieto dáta nemusia byť rovnakého dátového typu. Slovné premenné ako pohlavie, meno respondenta a podobne sú typu **character**. Vek respondenta je typu **integer**, hmotnosť, výška a ďalšie spojité premenné sú typu **double**. Preto v **R** existuje dátová štruktúra **data frame**, ktorá na prvý pohľad vyzerá ako matica, je to dvojrozmerný objekt, ktorý tvoria riadky a stĺpce a ktorý umožňuje, aby sa stĺpce mohli od seba odlišovať na základe dátového typu. V rámci jedného stĺpca však hodnoty musia mať rovnaký dátový typ. V reálnych analýzach totiž v stĺpcoch máme jednotlivé premenné a aj keď premenné sú rôzne, niektoré sú diskkrétne, iné spojité a podobne, hodnoty v rámci jednej premennej sú stále rovnakého typu (napríklad hmotnosť je stále spojitá premenná a jej hodnoty sú stále desatinné čísla). V objektoch typu **data frame** teda riadky predstavujú jednotlivé pozorovania a stĺpce jednotlivé premenné. Každý stĺpec môže obsahovať premenné iného typu a **data frame** tak vlastne predstavuje atomické vektory rôzneho typu spojené vedľa seba po stĺpcoch.

Objekt typu **data frame** vytvárame pomocou funkcie **data.frame()**, pričom vytváranie objektu typu **data frame** je veľmi podobné vytváraniu objektu typu **list**:

```

1 > novyDataFrame = data.frame(
2 +   "meno" = c("Jana", "Jakub", "Tereza"),
3 +   "vek" = c(28, 42, 19),
4 +   "hmotnost" = c(66.7, 91.5, 56.8),
5 +   "surodenci" = c(TRUE, FALSE, TRUE),
6 +   "pohlavie" = c("zena", "muz", "zena")
7 + )
8 >
9 > print(novyDataFrame)
10   meno vek hmotnost surodenci pohlavie
11 1   Jana  28    66.7      TRUE   zena
12 2  Jakub  42    91.5     FALSE   muz
13 3 Tereza  19    56.8      TRUE   zena
14 > str(novyDataFrame)
15 'data.frame': 3 obs. of  5 variables:

```

```

16 $ meno      : Factor w/ 3 levels "Jakub","Jana",...: 2 1 3
17 $ vek       : num  28 42 19
18 $ hmotnost  : num  66.7 91.5 56.8
19 $ surodenci: logi  TRUE FALSE TRUE
20 $ pohlavie  : Factor w/ 2 levels "muz","zena": 2 1 2
21 > is.matrix(novyDataFrame)
22 [1] FALSE
23 > is.data.frame(novyDataFrame)
24 [1] TRUE

```

Objekt typu data frame je v mnohom veľmi podobný matici. Nasledujúci zoznam uvádza niektoré z funkcií, ktorých použitie sme si už vysvetľovali v Kapitole 6.1.2 a ktoré majú rovnaké použitie aj pri práci s objektom typu data frame:

- `nrow()`: zistenie alebo určenie počtu riadkov,
- `ncol()`: zistenie alebo určenie počtu stĺpcov,
- `rownames()`: práca s názvami riadkov,
- `colnames()`: práca s názvami stĺpcov,
- `dim()`: zistenie rozmeru (počet riadkov a počet stĺpcov),
- `length()`: zistenie počtu hodnôt (počet riadkov vynásobený počtom stĺpcov),
- `head()`: vypísanie hodnôt z prvých pár riadkov,
- `tail()`: vypísanie hodnôt z posledných pár riadkov,
- `summary()`: súhrnné štatistiky pre všetky stĺpce,
- namiesto funkcie `rbind()` pri tvorbe objektu typu data frame riadkovým spájaním vektorov používame funkciu `rbind.data.frame()`,
- namiesto funkcie `cbind()` pri tvorbe objektu typu data frame stĺpcovým spájaním vektorov používame funkciu `cbind.data.frame()`.

Ak chceme pracovať s niektorými hodnotami z konkrétnych riadkov alebo stĺpcov, môžeme pri tom používať odkazovanie sa na poradie riadka a stĺpca ako v prípade matice. Tým, že ide o nehomogénnu dátovú štruktúru, tak v prípade objektu typu data frame pri práci s jednotlivými stĺpcami môžeme používať aj `$`, za ktorým nasleduje menovka stĺpca, s ktorým chceme pracovať:

```

1 > print(novyDataFrame)
2   meno vek hmotnost surodenci pohlavie
3 1  Jana  28    66.7      TRUE    zena
4 2  Jakub 42    91.5     FALSE    muz

```

```

5 3 Tereza 19 56.8 TRUE zena
6 > novyDataFrame[, 3]
7 [1] 66.7 91.5 56.8
8 > novyDataFrame$hmotnost
9 [1] 66.7 91.5 56.8
10 > novyDataFrame[-1, ]
11      meno vek hmotnost surodenci pohlavie
12 2 Jakub 42 91.5 FALSE muz
13 3 Tereza 19 56.8 TRUE zena
14 > novyDataFrame[, -4]
15      meno vek hmotnost pohlavie
16 1 Jana 28 66.7 zena
17 2 Jakub 42 91.5 muz
18 3 Tereza 19 56.8 zena

```

Vzhľadom na to, že matica a data frame sú na prvý pohľad veľmi podobné dátové štruktúry, ktoré sa od seba odlišujú iba na základe homogenity hodnôt, existujú funkcie, ktoré umožnia zmeniť maticu na data frame a naopak. Konverzia matice na data frame je jednoduchšia a nedochádza pri nej k strate informácie, nakoľko tam nedochádza k úprave hodnôt na nadradený dátový typ. Takúto konverziu vieme vykonať pomocou funkcie `as.data.frame()`. Naopak, pri konverzii objektu typu data frame na maticu sa vplyvom atomickosti môže stať, že všetky hodnoty budú napríklad typu `character`: ak v pôvodnom data frame bol stĺpec so slovnými hodnotami. Takto sa správa funkcia `as.matrix()`. Okrem tejto funkcie existuje ešte funkcia `data.matrix()`, ktorá všetky hodnoty objektu typu data frame najprv zmení na reálne čísla a tie potom uloží do matice.

```

1 > print(mat.RRR)
2      A      B      C      D      E      F
3 a 31.00 79.00 51.00 14.00 67.00 42.00
4 b 50.00 43.00 14.00 25.00 90.00 91.00
5 c 69.00 91.00 57.00 92.00 9.00 93.00
6 d 99.00 72.00 26.00 7.00 42.00 9.00
7 e 0.29 0.46 0.33 0.59 0.76 0.14
8 f 0.79 0.96 0.95 0.29 0.22 0.23
9 g 0.41 0.45 0.89 0.15 0.32 0.47
10 h 0.88 0.68 0.69 0.96 0.23 0.27
11 i 0.94 0.57 0.64 0.90 0.14 0.86
12 j 0.05 0.10 0.99 0.69 0.41 0.05
13 k 0.53 0.90 0.66 0.80 0.41 0.44
14 l 0.89 0.25 0.71 0.02 0.37 0.80
15 m 0.55 0.04 0.54 0.48 0.15 0.12
16 > is.matrix(mat.RRR)
17 [1] TRUE
18 > is.data.frame(mat.RRR)
19 [1] FALSE
20 > mat.RRR$C

```

```

21 Error in mat.RRR$C : $ operator is invalid for atomic vectors
22 > df.RRR <- as.data.frame(mat.RRR)
23 > is.matrix(df.RRR)
24 [1] FALSE
25 > is.data.frame(df.RRR)
26 [1] TRUE
27 > df.RRR$C
28 [1] 51.00 14.00 57.00 26.00 0.33 0.95 0.89 0.69 0.64 0.99
29 [11] 0.66 0.71 0.54

```

```

1 > print(novyDataFrame)
2   meno vek hmotnost surodenci pohlavie
3 1   Jana  28    66.7      TRUE    zena
4 2  Jakub  42    91.5     FALSE    muz
5 3 Tereza  19    56.8      TRUE    zena
6 > atomickaMatica <- as.matrix(novyDataFrame)
7 > print(atomickaMatica)
8   meno      vek  hmotnost surodenci pohlavie
9 [1,] "Jana"    "28"  "66.7"   "TRUE"   "zena"
10 [2,] "Jakub"   "42"  "91.5"   "FALSE"  "muz"
11 [3,] "Tereza"  "19"  "56.8"   "TRUE"   "zena"
12 > is.data.frame(atomickaMatica)
13 [1] FALSE
14 > is.matrix(atomickaMatica)
15 [1] TRUE
16 >
17 > print(novyDataFrame)
18   meno vek hmotnost surodenci pohlavie
19 1   Jana  28    66.7      TRUE    zena
20 2  Jakub  42    91.5     FALSE    muz
21 3 Tereza  19    56.8      TRUE    zena
22 > maticaRealnychCisel <- data.matrix(novyDataFrame)
23 > print(maticaRealnychCisel)
24   meno vek hmotnost surodenci pohlavie
25 [1,]  2  28    66.7      1      2
26 [2,]  1  42    91.5      0      1
27 [3,]  3  19    56.8      1      2
28 > is.data.frame(maticaRealnychCisel)
29 [1] FALSE
30 > is.matrix(maticaRealnychCisel)
31 [1] TRUE

```

Komplexnejšiu prácu s objektmi typu data frame uvádzame v Kapitole 8.2

6.3 Zistenie obsahu dátovej štruktúry

Ak chceme zistiť štruktúru objektu uloženého v **R**, využívame pri tom funciu `str()`, pomocou ktorej je vypísaná trieda premennej a taktiež sa zobrazí jej štruktúra (`rownames`, `colnames`...).

```
1 > print(vektor1)
2  A  B  C  D  E  F  G  H  I  J
3  1  2  3  4  5  6  7  8  9 10
4 > str(vektor1)
5  Named int [1:10] 1 2 3 4 5 6 7 8 9 10
6  - attr(*, "names")= chr [1:10] "A" "B" "C" "D" ...
7 >
8 > print(mat.A)
9      [,1] [,2] [,3] [,4] [,5] [,6]
10 [1,]   15   11    7    3   -1   -5
11 [2,]   14   10    6    2   -2   -6
12 [3,]   13    9    5    1   -3   -7
13 [4,]   12    8    4    0   -4   -8
14 > str(mat.A)
15  num [1:4, 1:6] 15 14 13 12 11 10 9 8 7 6 ...
16 >
17 > print(novyList)
18 $hmotnost
19 [1] 66.7 91.5 56.8
20
21 $surodenci
22 [1] TRUE FALSE TRUE
23
24 $muz
25 [1] FALSE TRUE FALSE
26
27 > str(novyList)
28 List of 3
29 $ hmotnost : num [1:3] 66.7 91.5 56.8
30 $ surodenci: logi [1:3] TRUE FALSE TRUE
31 $ muz      : logi [1:3] FALSE TRUE FALSE
32 >
33 > print(novyDataFrame)
34      meno vek hmotnost surodenci pohlavie
35 1   Jana  28    66.7      TRUE    zena
36 2  Jakub  42    91.5     FALSE    muz
37 3 Tereza  19    56.8      TRUE    zena
38 > str(novyDataFrame)
39 'data.frame': 3 obs. of  5 variables: #obs. = pocet pozorovani
      (observations), variables = pocet premennych
```

```

40 $ meno      : Factor w/ 3 levels "Jakub","Jana",...: 2 1 3
41 $ vek       : num  28 42 19
42 $ hmotnost  : num  66.7 91.5 56.8
43 $ surodenci: logi  TRUE FALSE TRUE
44 $ pohlavie  : Factor w/ 2 levels "muz","zena": 2 1 2
45 > print(vektor1)
46  A  B  C  D  E  F  G  H  I  J
47  1  2  3  4  5  6  7  8  9 10
48 > str(vektor1)
49  Named int [1:10] 1 2 3 4 5 6 7 8 9 10
50 - attr(*, "names")= chr [1:10] "A" "B" "C" "D" ...
51 >
52 > print(mat.A)
53      [,1] [,2] [,3] [,4] [,5] [,6]
54 [1,]   15   11    7    3   -1   -5
55 [2,]   14   10    6    2   -2   -6
56 [3,]   13    9    5    1   -3   -7
57 [4,]   12    8    4    0   -4   -8
58 > str(mat.A)
59  num [1:4, 1:6] 15 14 13 12 11 10 9 8 7 6 ...
60 >
61 > print(novyList)
62 $hmotnost
63 [1] 66.7 91.5 56.8
64
65 $surodenci
66 [1]  TRUE FALSE  TRUE
67
68 $muz
69 [1] FALSE  TRUE FALSE
70
71 > str(novyList)
72 List of 3
73 $ hmotnost : num [1:3] 66.7 91.5 56.8
74 $ surodenci: logi [1:3] TRUE FALSE TRUE
75 $ muz      : logi [1:3] FALSE TRUE FALSE
76 >
77 > print(novyDataFrame)
78      meno vek hmotnost surodenci pohlavie
79 1   Jana  28    66.7      TRUE    zena
80 2  Jakub 42    91.5     FALSE    muz
81 3 Tereza 19    56.8      TRUE    zena
82 > str(novyDataFrame)
83 'data.frame': 3 obs. of  5 variables:
84 $ meno      : Factor w/ 3 levels "Jakub","Jana",...: 2 1 3
85 $ vek       : num  28 42 19

```

```
86 $ hmotnost : num 66.7 91.5 56.8
87 $ surodenci: logi TRUE FALSE TRUE
88 $ pohlavie : Factor w/ 2 levels "muz","zena": 2 1 2
```

Kapitola 7

Základy spracovania dát

Doteraz sme si stále vytvárali objekty a rôzne ilustračné dátové štruktúry priamo v **R**. V praxi však najčastejšie potrebujeme spracovať údaje, ktoré sú uložené v nejakom externom súbore (napríklad Excelovská tabuľka) a potrebujeme tieto dáta importovať. Cieľom kapitoly je ukázať, ako správne načítať dáta do **R**, ako ich správne spracovať a taktiež ako ich následne uložiť do rôznych formátov. Hneď v úvode treba povedať, že programovací jazyk **R** je v tomto smere veľmi flexibilný a okrem štandardných dátových formátov typu (.txt, .csv, .xls alebo .xlsx) dokáže pracovať aj s inými formátmi (napríklad databázy komerčných štatistických programov SPSS, Stata, systat či SAS).

7.1 Import dát

Samotnému načítaniu údajov do **R** by mala predchádzať príprava súboru na prácu v tomto programovacom jazyku, preto odporúčame otvoriť si dátový súbor a riadiť sa nasledovnými pravidlami:

- Mená riadkov a stĺpcov:
 - Prvý riadok používajte na pomenovanie stĺpcov, stĺpce obvykle reprezentujú jednotlivé premenné.
 - Prvý stĺpec používajte na pomenovanie riadkov, riadky obvykle reprezentujú jednotlivé pozorovania.
 - Každé pozorovanie musí mať jedinečné pomenovanie, preto ak chcete mať riadky pomenované, skontrolujte si, či sa niektoré pomenovania riadkov neopakujú, lebo inak budete konfrontovaní s nasledujúcou chybovou hláškou:

```
1 Error in read.table(file = file, header = header, sep = sep
  , quote = quote, :
```

- Samozrejme, aj každá premenná by mala mať jedinečné pomenovanie. Ak by sme však v dátovom súbore mali viacero stĺpcov s rovnakými názvami, **R** dátový súbor načíta, pričom stĺpce, ktoré majú v zdrojovom súbore rovnaké názvy, bude premenovávať prostredníctvom bodky, za ktorou nasleduje číslo v poradí od jednotky nahor. Napríklad ak máme viacero stĺpcov pomenovaných ako `variable`, potom prvý v poradí bude mať menovku `variable`, druhý bude mať menovku `variable.1`, tretí bude mať menovku `variable.2` a podobne.
- Pravidlá pre pomenovanie stĺpcov:
 - Pri pomenovaní stĺpcov sa vyvarujte používaniu interpunkčných znamienok, ktoré by pri importe a následnom spracovaní údajov mohli byť zdrojom zbytočných chýb.
 Neodporúčaný názov stĺpca: Štatistika
 Odporúčaný názov stĺpca: Statistika
 - Rovnako by menovky stĺpcov nemali obsahovať medzery. Pri načítaní sú tieto medzery nahradené bodkou, na to je potom potrebné myslieť aj pri ich následnom spracovaní.
 Neodporúčaný názov stĺpca: Pouzivaj R
 Odporúčané názvy stĺpca: Pouzivaj.R, Pouzivaj_R prípadne PouzivajR
 - Nepoužívajte na pomenovanie stĺpcov špeciálne symboly ako napríklad: `?`, `$`, `*`, `+`, `#`, `(`, `)`, `-`, `/`, `}`, `{`, `|`, `>`, `<` a podobne, ktoré majú v programe **R** osobitnú funkcionálnu hodnotu. Ako sme uviedli v predchádzajúcom bode, použitie podčiarkovníka (`_`) je povolené.
 - Pomenovanie stĺpcov by nemalo začínať číslicou, ale písmenom, v opačnom prípade by to mohlo opäť byť zdrojom zbytočných chýb.
 Neodporúčaný názov stĺpca: 4ST110
 Odporúčaný názov stĺpca: predmet_4ST110
 - Aj pri pomenovaní stĺpcov platí, že **R** odlišuje veľké a malé písmená (case sensitivity).
- Ďalšie pravidlá pre načítavanie dát:
 - Skontrolujte, či sa v súbore nenachádzajú prázdne riadky, ak áno, vymažte ich.
 - Ak v zdrojovom dátovom súbore máte akékoľvek poznámky alebo komentáre, vymažte ich.
 - Ak v zdrojovom súbore máte nejaké chýbajúce hodnoty (prázdne oblasti v tabuľke) nahraďte ich dvoma písmenami `NA` (not available: nedostupné, chýbajúce hodnoty).

-
- Pri práci s dátumami používate štvorciferný zápis roku (four digit format).
Neodporúčaný formát dátumu: 01/01/20
Odporúčaný formát dátumu: 01/01/2020
 - Zdrojový súbor, ktorý následne importujeme do **R**, odporúčame uložiť buď vo formáte textového súboru oddeleného tabulátormi `.txt` (tab-delimited text file), alebo ako súbor s hodnotami oddelených čiarkou `.csv` (comma separated value file).

7.1.1 Funkcie vhodné na import dát a ich základné parametre

Všeobecnou funkciou slúžiacou na načítanie údajov v tabuľkovom formáte do **R** je funkcia `read.table()`. Údaje sa načítajú ako data frame. Najdôležitejšie argumenty tejto funkcie sú nasledovné:

- **file**: Meno súboru aj s príponou v úvodzovkách, z ktorého majú byť načítané údaje. Každý riadok tabuľky je načítaný z príslušných riadkov zdrojovej tabuľky. Ak je uvedený iba názov súboru s príponou, súbor sa načítava z pracovného adresára:

```
1 > read.table("data.csv")
```

Pre viac informácií o nastavení a zmene pracovného adresára viď Kapitulu 1.9. V tomto argumente môže byť uvedená aj absolútna adresa zdrojového súboru, v tom prípade je načítanie súboru nasledovné:

```
1 > read.table("D:/Users/Files/data.csv")
```

Alternatívou môže byť využitie funkcie `file.choose()`, ktorá otvorí nové dialógové okno, prostredníctvom ktorého je možné súbor jednoducho vyhľadať v počítači:

```
1 > read.table(file.choose())
```

- **header**: Logická hodnota, pri ktorej definujeme, či sa v prvom riadku nachádzajú menovky stĺpcov. Tento argument je prednastavený na `FALSE`, ak však načítavaný dátový súbor v prvom riadku obsahuje o jednu hodnotu menej, ako je počet stĺpcov, automaticky sa nastaví na `TRUE`.
- **sep**: Ide o argument, ktorým definujeme, ako sú oddelené hodnoty v načítavacom zdrojovom súbore. Vo východiskovom tvare je nastavený pre túto funkciu ako `sep = ""`, čo znamená, že na oddelenie hodnôt zdrojového súboru boli použité „prázdne miesta“ (jedna alebo viac medzier, tabulátor, nový riadok a podobne).
- **dec**: Oddeľovač desatinných miest typu `character`, väčšinou ide o desatinnú čiarku (typické pre stredoeurópske databázy) a desatinnú bodku (typické pre americké databázy).

- **row.names**: Vektor riadkových názvov (menoviek riadkov), môže ísť o vektor, ktorým tieto názvy nadefinujeme alebo to môže byť celé číslo z intervalu 1 až m , kde m je počet stĺpcov načítavaného dátového súboru. V tom prípade poradím stĺpca určíme, z ktorého stĺpca sa majú načítavať menovky riadkov. Najčastejšie sa menovky riadkov nachádzajú v prvom stĺpci, preto tento argument väčšinou vyzerá nasledovne: **row.names** = 1. Ak máme nadefinovaný argument **header** (menovky stĺpcov) a prvý riadok obsahuje o jednu hodnotu menej, ako je počet stĺpcov, potom sa automaticky za menovky riadkov nastaví prvý stĺpec. Inak, ak tento argument abscentuje, tak riadky sa pomenujú číslaním. V prípade, ak sa tento argument nastaví ako **row.names** = NULL, automaticky sa riadky pomenujú číslaním.
- **col.names**: Vektor s menovkami pre premenné v jednotlivých stĺpcoch. Pokiaľ argument **header** = FALSE, automaticky sa označujú ako V1, V2... podľa poradia stĺpca.
- **nrows**: Celé číslo predstavujúce maximálny počet riadkov, ktoré sa majú načítať.
- **skip**: Vyjadruje počet riadkov, ktoré sa majú na začiatku preskočiť.
- **quote**: Textový reťazec typu **character** obsahujúci znak, ktorý sa používa ako úvodzovky k označeniu textu.
- **comment.char**: Textový reťazec typu **character** obsahujúci znak, ktorý sa používa k označeniu komentárov.
- **na.string**: Vektor obsahujúci hodnoty, ktoré sa majú nahradiť hodnotou NA.
- **colClasses**: Vektor, ktorý určí, aký dátový typ majú jednotlivé stĺpce.

Na uľahčenie práce, aby nebolo potrebné definovať stále pri načítavaní niektoré argumenty funkcie **read.table()**, existujú funkcie, ktoré majú automaticky nadefinované niektoré argumenty podľa toho, s akou databázou pracujeme. Vo všeobecnosti existujú dva typy databáz: prvá je typická pre európske databázy, ktoré na oddeľovanie desatinných miest väčšinou používajú desatinnú čiarku (,) a bodkočiarku (v ČJ středník ;) používajú na oddeľovanie hodnôt. Americké databázy naopak používajú na oddeľovanie desatinných miest bodku (v ČJ tečku .) a čiarku (,) používajú na oddeľovanie hodnôt:

- **read.csv()**: Na načítanie súborov s hodnotami oddelenými čiarkou (.csv) v amerických databázach. Východiskové nastavenie funkcie je nasledovné:

```
1 > read.csv(file, header = TRUE, sep = ",", dec = ".", ...)
```

- **read.csv2()**: Na načítanie súborov s hodnotami oddelenými čiarkou (.csv) v európskych databázach. Východiskové nastavenie funkcie je nasledovné:

```
1 > read.csv2(file, header = TRUE, sep = ";", dec = ",", ...)
```

- `read.delim()`: Na načítanie súborov s hodnotami oddelenými tabulátormi (.txt) v amerických databázach. Východiskové nastavenie funkcie je nasledovné:

```
1 > read.delim(file, header = TRUE, sep = "\t", dec = ".", ...)
```

- `read.delim2()`: Na načítanie súborov s hodnotami oddelenými tabulátormi (.txt) v európskych databázach. Východiskové nastavenie funkcie je nasledovné:

```
1 > read.delim2(file, header = TRUE, sep = "\t", dec = ",", ...)
```

V prostredí RStudio je v okne pracovného prostredia (Environment) tlačidlo **Import Dataset**, pomocou ktorého veľmi jednoduchým klikacím spôsobom môžeme vybrať zdrojové údaje, ktoré chceme načítať.

7.1.2 Načítavanie dát z internetu

Ak je zdrojový súbor uložený na nejakej internetovej adrese po kliknutí, na ktorú sa súbor hneď zobrazí, môžeme ho načítať priamo bez toho, aby sme ho museli ukladať a načítať z pracovného adresára. Môžeme pri tom využiť všetky vyššie uvedené funkcie v závislosti od toho, o aký typ databázy ide. Predtým, ako vykonáme načítanie údajov, pozrime sa, ako súbor vyzerá, uvažujme pri tom vzorový dátový súbor rozdelenia domácich prác medzi manželku a manžela: <http://www.sthda.com/sthda/RDoc/data/housetasks.txt>.

```
1 > zdroj <- "http://www.sthda.com/sthda/RDoc/data/housetasks.txt"
2 > data01 <- read.delim(file = zdroj)
3 > print(data01)
```

		X	Wife	Alternating	Husband	Jointly
1	Laundry	156	14	2	4	
2	Main_meal	124	20	5	4	
3	Dinner	77	11	7	13	
4	Breakfeast	82	36	15	7	
5	Tidying	53	11	1	57	
6	Dishes	32	24	4	53	
7	Shopping	33	23	9	55	
8	Official	12	46	23	15	
9	Driving	10	51	75	3	
10	Finances	13	13	21	66	
11	Insurance	8	1	53	77	
12	Repairs	0	3	160	2	
13	Holidays	0	1	6	153	

Vytvoríme nový objekt, v ktorom rovnakú databázu načítame s tým, že chceme, aby hodnoty v prvom stĺpci `row.names = 1` boli považované za menovky riadkov:

```

1 > data02 <- read.delim(file = zdroj, row.names = 1)
2 > print(data02)
3      Wife Alternating Husband Jointly
4 Laundry      156         14        2        4
5 Main_meal    124         20        5        4
6 Dinner       77         11        7       13
7 Breakfast    82         36       15        7
8 Tidying      53         11        1       57
9 Dishes       32         24        4       53
10 Shopping    33         23        9       55
11 Official    12         46       23       15
12 Driving     10         51       75        3
13 Finances    13         13       21       66
14 Insurance     8          1       53       77
15 Repairs      0          3      160        2
16 Holidays    0          1        6      153

```

7.1.3 Načítavanie dát z pracovného adresára

Na príklade vybraných dátových súborov si ukážeme načítavanie údajov prostredníctvom vyššie uvedených funkcií. Údaje je potrebné stiahnuť a uložiť do pracovného adresára, na čo využijeme funkciu `download.file()`.

Súbor .txt

Prostredníctvom argumentu `destfile = file.path(getwd(), "textovyInput.txt")` nastavíme, aby sa súbor, ktorý chceme stiahnuť, volal `textovyInput.txt` a uložil sa do aktuálneho pracovného adresára. Pre viac informácií o nastavení a zmene pracovného adresára pozri Kapitolu 1.9. Ak má náš pracovný adresár napríklad adresu `D:/Users`, môžeme si ju vypísať prostredníctvom funkcie `getwd()`:

```

1 > getwd()
2 [1] "D:/Users"

```

```

1 > download.file(
2 +   "https://rstatistika.vse.cz/wp-content/uploads/page/174/mtcars.
3 +   txt",
4 +   destfile = file.path(getwd(), "textovyInput.txt")
5 + )
6 trying URL 'https://rstatistika.vse.cz/wp-content/uploads/page/174/
7   mtcars.txt'
8 Content type 'text/plain' length 1750 bytes
9 downloaded 1750 bytes

```

Keď si následne otvoríme stiahnutý súbor `textovyInput.txt`, vidíme, že ide o americký typ databázy, preto na jej načítanie využijeme funkciu `read.csv()`. Vzhľadom na to, že načítaný dátový súbor obsahuje vyše 30 riadkov, z dôvodu prehľadnosti si ho nebudeme zobrazovať celý, ale uvedieme iba jeho pár prvých riadkov (funkcia `head()`):

```
1 > data03 <- read.csv(file = "textovyInput.txt")
2 > head(data03)
3           X mpg cyl disp  hp drat   wt  qsec vs am gear
4 1      Mazda RX4 21.0   6  160 110 3.90 2.620 16.46 0  1   4
5 2    Mazda RX4 Wag 21.0   6  160 110 3.90 2.875 17.02 0  1   4
6 3      Datsun 710 22.8   4  108  93 3.85 2.320 18.61 1  1   4
7 4    Hornet 4 Drive 21.4   6  258 110 3.08 3.215 19.44 1  0   3
8 5 Hornet Sportabout 18.7   8  360 175 3.15 3.440 17.02 0  0   3
9 6      Valiant 18.1   6  225 105 2.76 3.460 20.22 1  0   3
10 carb
11 1     4
12 2     4
13 3     1
14 4     1
15 5     2
16 6     1
```

Vytvoríme nový objekt, v ktorom rovnakú databázu načítame s tým, že chceme, aby hodnoty v prvom stĺpci `row.names = 1` boli považované za menovky riadkov:

```
1 > data04 <- read.csv(file = "textovyInput.txt", row.names = 1)
2 > head(data04)
3           mpg cyl disp  hp drat   wt  qsec vs am gear carb
4 Mazda RX4    21.0   6  160 110 3.90 2.620 16.46 0  1   4   4
5 Mazda RX4 Wag 21.0   6  160 110 3.90 2.875 17.02 0  1   4   4
6 Datsun 710    22.8   4  108  93 3.85 2.320 18.61 1  1   4   1
7 Hornet 4 Drive 21.4   6  258 110 3.08 3.215 19.44 1  0   3   1
8 Hornet Sportabout 18.7   8  360 175 3.15 3.440 17.02 0  0   3   2
9 Valiant      18.1   6  225 105 2.76 3.460 20.22 1  0   3   1
```

Ako sme uviedli vyššie, funkcie `read.csv()`, respektíve `read.csv2()` môžeme používať aj v prípade, ak chceme načítať súbor vo formáte `.txt`. Ak by sme pri jeho načítaní použili funkcie `read.delim()` alebo `read.table()`, je potrebné dodefinovať argument `sep`, nakoľko pri oboch týchto funkciách je definovaný na „prázdne miesta“ (medzery tabulátory a podobne), avšak v zdrojovom súbore nemáme v riadkoch medzery, ale hodnoty sú oddelené čiarkou (`,`). Všimnime si, že v tom prípade načíta všetky hodnoty v rámci riadka ako jeden textový reťazec, preto je výsledný rozmer vytvoreného objektu typu `data frame` $n \times 1$, kde n je počet pozorovaní:

```
1 > data05 <- read.delim(file = "textovyInput.txt")
2 > head(data05)
```



```

3           X.mpg.cyl.disp.hp.drat.wt.qsec.vs.am.gear.carb
4 1           Mazda RX4,21,6,160,110,3.9,2.62,16.46,0,1,4,4
5 2           Mazda RX4 Wag,21,6,160,110,3.9,2.875,17.02,0,1,4,4
6 3           Datsun 710,22.8,4,108,93,3.85,2.32,18.61,1,1,4,1
7 4   Hornet 4 Drive,21.4,6,258,110,3.08,3.215,19.44,1,0,3,1
8 5   Hornet Sportabout,18.7,8,360,175,3.15,3.44,17.02,0,0,3,2
9 6           Valiant,18.1,6,225,105,2.76,3.46,20.22,1,0,3,1
10 > dim(data05)
11 [1] 32  1

```

Zadajme argument pre oddeľovanie hodnôt (`sep = ","`), pričom chceme, aby hodnoty v prvom stĺpci boli považované za menovky riadkov (`row.names = 1`):

```

1 > data06 <- read.delim(file = "textovyInput.txt", row.names = 1, sep
2   = ",")
3 > head(data06)
4           mpg cyl disp  hp drat   wt  qsec vs am gear carb
5 Mazda RX4    21.0   6  160 110 3.90 2.620 16.46 0  1   4    4
6 Mazda RX4 Wag 21.0   6  160 110 3.90 2.875 17.02 0  1   4    4
7 Datsun 710    22.8   4  108  93 3.85 2.320 18.61 1  1   4    1
8 Hornet 4 Drive 21.4   6  258 110 3.08 3.215 19.44 1  0   3    1
9 Hornet Sportabout 18.7   8  360 175 3.15 3.440 17.02 0  0   3    2
10 Valiant      18.1   6  225 105 2.76 3.460 20.22 1  0   3    1
11 > dim(data06)
12 [1] 32 11

```

Súbor .csv európskeho typu

Najprv si do nášho pracovného adresára stiahneme súbor `europskeCSV.csv`.

```

1 > download.file(
2 +   "https://rstatistika.vse.cz/wp-content/uploads/page/174/mtcarsA.
3   csv",
4 +   destfile = file.path(getwd(), "europskeCSV.csv")
5 + )
6 trying URL 'https://rstatistika.vse.cz/wp-content/uploads/page/174/
7   mtcarsA.csv'
8 Content type 'text/x-comma-separated-values' length 1816 bytes
9 downloaded 1816 bytes

```

Keď si následne otvoríme stiahnutý súbor `europskeCSV.csv`, vidíme, že ide o európsky typ databázy, preto na jej načítanie využijeme funkciu `read.csv2()`. Vzhľadom na to, že načítaný dátový súbor obsahuje vyše 30 riadkov, z dôvodu prehľadnosti si ho nebudeme zobrazovať celý, ale uvedieme iba jeho pár prvých riadkov (funkcia `head()`):

```

1 > data07 <- read.csv2(file = "europskeCSV.csv", row.names = 1)
2 > head(data07)
3           mpg  cyl  disp  hp  drat    wt    qsec vs  am gear carb
4 Mazda RX4      21.0   6  160  110  3.90  2.620  16.46  0   1    4    4
5 Mazda RX4 Wag  21.0   6  160  110  3.90  2.875  17.02  0   1    4    4
6 Datsun 710     22.8   4  108   93  3.85  2.320  18.61  1   1    4    1
7 Hornet 4 Drive  21.4   6  258  110  3.08  3.215  19.44  1   0    3    1
8 Hornet Sportabout 18.7   8  360  175  3.15  3.440  17.02  0   0    3    2
9 Valiant        18.1   6  225  105  2.76  3.460  20.22  1   0    3    1

```

Súbor .csv amerického typu

Najprv si do nášho pracovného adresára stiahneme súbor `americkeCSV.csv`.

```

1 > download.file(
2 +   "https://rstatistika.vse.cz/wp-content/uploads/page/174/mtcarsB.
3 +   csv",
4 +   destfile = file.path(getwd(), "americkeCSV.csv")
5 + )
6 trying URL 'https://rstatistika.vse.cz/wp-content/uploads/page/174/
7   mtcarsB.csv'
8 Content type 'text/x-comma-separated-values' length 1816 bytes
9 downloaded 1816 bytes

```

Keď si následne otvoríme stiahnutý súbor `americkeCSV.csv`, vidíme, že ide o americký typ databázy, preto na jej načítanie využijeme funkciu `read.csv()`. Vzhľadom na to, že načítaný dátový súbor obsahuje vyše 30 riadkov, z dôvodu prehľadnosti si ho nebudeme zobrazovať celý, ale uvedieme iba jeho pár prvých riadkov (funkcia `head()`):

```

1 > data08 <- read.csv(file = "americkeCSV.csv", row.names = 1)
2 > head(data08)
3           mpg  cyl  disp  hp  drat    wt    qsec vs  am gear carb
4 Mazda RX4      21.0   6  160  110  3.90  2.620  16.46  0   1    4    4
5 Mazda RX4 Wag  21.0   6  160  110  3.90  2.875  17.02  0   1    4    4
6 Datsun 710     22.8   4  108   93  3.85  2.320  18.61  1   1    4    1
7 Hornet 4 Drive  21.4   6  258  110  3.08  3.215  19.44  1   0    3    1
8 Hornet Sportabout 18.7   8  360  175  3.15  3.440  17.02  0   0    3    2
9 Valiant        18.1   6  225  105  2.76  3.460  20.22  1   0    3    1

```

Excelovské súbory .xls a .xlsx

Najprv si do nášho pracovného adresára stiahneme súbor `excelovskyInput.xlsx`. Ak chceme načítavať Excelovské súbory, je potrebné pridať ešte jeden argument (`mode =`

"wb") funkcie `download.file()`. Argument udáva mód, v akom budeme so súborom pracovať (wb predstavuje binárny súbor pre zápis). Východiskové nastavenie je `mode = "w"` (w predstavuje textový súbor pre zápis), čo by nám síce súbor stiahlo, ale mohli by sme mať problém s jeho zobrazením.

```
1 > download.file(  
2 +   "https://rstatistika.vse.cz/wp-content/uploads/page/174/data.  
   xlsx",  
3 +   destfile = file.path(getwd(), "excelovskyInput.xlsx"),  
4 +   mode = "wb"  
5 + )  
6 trying URL 'https://rstatistika.vse.cz/wp-content/uploads/page/174/  
   data.xlsx'  
7 Content type 'application/vnd.openxmlformats-officedocument.  
   spreadsheetml.sheet' length 12984 bytes (12 KB)  
8 downloaded 12 KB
```

Keď si následne otvoríme stiahnutý súbor `excelovskyInput.xlsx`, vidíme, že obsahuje dva hárk (listy). Prvý hárok (list) je pomenovaný ako `mtcars` a druhý ako `regression`. Ak chceme načítavať excelovské súbory, najčastejšie pri tom používame balíky `readxl`, `xlsx` a `openxlsx`.

Výhodou balíka `readxl` je, že dokáže načítavať ako starší formát `.xls` (zošit Excelu 1997–2003) tak novší formát `.xlsx`. Pre následné uloženie údajov odporúčame využívať balíky `xlsx` a `openxlsx`.

Pre načítanie Excelovského súboru môžeme použiť funkciu `read_excel()`, ktorá má podobné argumenty ako funkcia `read.table()`. Keďže máme viacero hárkov (v ČJ listov), je potrebné zadať aj názov alebo poradie háрку, ktorý chceme načítať. Na to slúži argument `sheet`.

Ak by sme chceli načítať druhý hárok (list) zo súboru `excelovskyInput.xlsx`, použijeme argument `sheet = 2` alebo `sheet = "regression"`, keďže druhý hárok má pomenovanie `regression`. Ak by sme argument nezadali, funkcia automaticky pracuje stále s dátovým súborom v prvom háрку (liste).

```
1 > if (! require ("readxl")) {  
2 +   install.packages ("readxl", dependencies = TRUE)  
3 +   library (readxl)  
4 + }  
5 >  
6 > data09 <- read_excel("excelovskyInput.xlsx", sheet = "regression")  
7 > print(data09)  
8 # A tibble: 23 x 2  
9       x         y  
10    <dbl> <dbl>  
11 1    12.4    11.2
```

```

12  2  14.3  12.5
13  3  14.5  12.7
14  4  14.9  13.1
15  5  16.1  14.1
16  6  16.9  14.8
17  7  16.5  14.4
18  8  15.4  13.4
19  9  17    14.9
20 10  17.9  15.6
21 # ... with 13 more rows

```

Všimnime si, že tentokrát objekt `data09` pomocou funkcie `read_excel()` nie je iba obyčajný data frame, ale ide o nejakú dátovú štruktúru typu `tibble`. Ide o data frame, v ktorom sú vylepšené niektoré funkcionality. Vzhľadom na to, že **R** je relatívne starší programovací jazyk, funkcionality, ktoré boli v minulosti užitočné, už teraz nie sú potrebné a, naopak, množstvo užitočných funkcií sa objavilo až po určitom čase. Keďže je zložitá meniť základy **R**, inovácie sa realizujú prostredníctvom balíkov. Užitočné rozšírenie ponúka balík `tidyverse`, ktorý pracuje s dátovými štruktúrami typu `tibble`. Balík `tidyverse` nebolo potrebné inštalovať, nakoľko sa nainštaloval a spustil pri inštalácii balíka `readxl`, keďže sme pri funkcii `install.packages()` nastavili argument `dependencies` na `TRUE`. Vzhľadom na komplexnosť danej problematiky sa v učebnici nebudeme štruktúre typu `tibble` hlbšie venovať. S danou štruktúrou však pracujeme podobne ako s objektom typu data frame:

```

1 > dim(data09)
2 [1] 23  2
3 > is.data.frame(data09)
4 [1] TRUE
5 > colnames(data09)
6 [1] "x" "y"
7 >
8 > data09$x
9 [1] 12.4 14.3 14.5 14.9 16.1 16.9 16.5 15.4 17.0 17.9 18.8 20.3
10 [13] 22.4 19.4 15.5 16.7 17.3 18.4 19.2 17.4 19.5 19.7 21.2
11 > data09$y
12 [1] 11.2 12.5 12.7 13.1 14.1 14.8 14.4 13.4 14.9 15.6 16.4 17.7
13 [13] 19.6 16.9 14.0 14.6 15.1 16.1 16.8 15.2 17.0 17.2 18.6
14 >
15 > colSums(data09)
16      x      y
17 401.7 351.9
18 > sum(data09[, 1])
19 [1] 401.7
20 > sum(data09[, 2])
21 [1] 351.9

```

```

22 >
23 > summary(data09)
24           x           y
25 Min.      :12.40   Min.      :11.20
26 1st Qu.:15.80   1st Qu.:14.05
27 Median :17.30   Median :15.10
28 Mean     :17.47   Mean     :15.30
29 3rd Qu.:19.30   3rd Qu.:16.85
30 Max.     :22.40   Max.     :19.60

```

7.2 Export dát z R do štandardných formátov

Rovnakým spôsobom, ako sme údaje načítavali, ich môžeme aj exportovať a uložiť do nami zvoleného formátu.

Funkcia `write.table()` je akousi inverznou funkciou k funkcii `read.table()`, keďže jedna slúži na export údajov a druhá na import. Rovnako fungujú aj funkcie `write.csv()` a `write.csv2()`. Všetky tieto funkcie majú navyše argument `x`, ktorý predstavuje maticu alebo data frame, ktorú chceme uložiť. Vytvoríme si v objekte s názvom `mocniny` data frame, ktorého prvý stĺpec bude obsahovať písmená latinskej abecedy, druhý stĺpec bude obsahovať čísla od 1 po počet písmen latinskej abecedy, tretí stĺpec bude obsahovať druhé mocniny hodnôt v druhom stĺpci a štvrtý stĺpec bude obsahovať tretie mocniny hodnôt v druhom stĺpci. Menovky riadkov tohto objektu nech budú písmená v prvom stĺpci:

```

1 > mocniny <- cbind.data.frame(LETTERS,1:length(LETTERS))
2 > colnames(mocniny) <- c("pismo", "cislo")
3 > head(mocniny)
4   pismo  cislo
5 1      A      1
6 2      B      2
7 3      C      3
8 4      D      4
9 5      E      5
10 6      F      6

```

Keďže ide o objekt typu data frame, môžeme vložiť stĺpce s druhými a tretími mocninami veľmi jednoducho:

```

1 > mocniny$naDruhu <- mocniny$cislo ^ 2
2 > mocniny$naTretiu <- mocniny$cislo ^ 3
3 > head(mocniny)
4   pismo  cislo  naDruhu  naTretiu
5 1      A      1         1         1
6 2      B      2         4         8

```

7	3	C	3	9	27
8	4	D	4	16	64
9	5	E	5	25	125
10	6	F	6	36	216

Ešte nám ostáva pomenovať riadky hodnotami v prvom stĺpci. Na to nám poslúži funkcia `rownames()`:

```

1 > rownames(mocniny) <- mocniny$pismo
2 > head(mocniny)
3   pismo  cislo naDruhu naTretiu
4 A      A      1       1        1
5 B      B      2       4        8
6 C      C      3       9       27
7 D      D      4      16       64
8 E      E      5      25      125
9 F      F      6      36      216

```

Mohlo by sa zdať, že vzhľadom na to, že máme v stĺpci `pismo` objektu `mocniny` uložené tie isté hodnoty ako v menovkách riadkov, je možné tento stĺpec vymazať. V skutočnosti je však bezpečnejšie uchovávať informácie o jednotlivých pozorovaniach v osobitnom stĺpci, nakoľko niektoré funkcie, ktoré sa používajú k spracovávaniu databáz, mená týchto riadkov odstraňujú. Rovnako tak aj menovky riadkov nepodporujú niektoré formáty, do ktorých sa dáta ukladajú, čím by sme mohli prísť o dôležité informácie.

Uložme teraz vytvorený objekt `mocniny` do pracovného adresára vo formáte `.txt` `.csv` a `.xlsx`:

Súbor `.txt`

```

1 > write.table(x = mocniny, file = "mocniny.txt", sep = ",")

```

Súbor `.csv` európskeho typu

```

1 > write.csv2(x = mocniny, file = "mocninyEU.csv")

```

Súbor `.csv` amerického typu

```

1 > write.csv(x = mocniny, file = "mocninyUSA.csv")

```

Súbor .xlsx

```
1 > if (!require ("openxlsx")) {  
2 +   install.packages ("openxlsx", dependencies = TRUE)  
3 +   library (openxlsx)  
4 + }  
5 >  
6 > write.xlsx(  
7 +   mocniny,  
8 +   file = "mocniny.xlsx" ,  
9 +   sheetName = "mocniny",  
10 +   col.names = TRUE,  
11 +   row.names = TRUE,  
12 + )
```

V tejto časti upozorňujeme na relatívne častú chybu, ktorá sa vyskytuje pri exporte údajov z **R** v ľubovoľnom formáte, ktorá je sprevádzaná týmto upozornením:

```
1 Warning message:  
2 In file.create(to[okay]) :  
3   cannot create file 'mocniny.xlsx', reason 'Permission denied'
```

Toto upozornenie sa vypisuje vtedy, ak súbor, ktorý chceme vytvoriť, máme súčasne otvorený, čo sa často stáva v prípade, ak súbor vytvárame na viackrát a jeho otváraním kontrolujeme, či má požadovaný tvar a formu. Ak súbor zatvoríme a skript zbehneme ešte raz, problém sa tým vyrieši.

7.3 Export dát z R do formátu RData

7.3.1 Funkcie saveRDS() a readRDS()

Ak chceme v prostredí **R** ukladať samotné objekty (nemusí ísť iba o databázy, ale vo všeobecnosti o akýkoľvek objekt), je možné na to využiť funkcie:

- `saveRDS()`,
- `save()`,
- `save.image()`.

Rozdiel medzi funkciami `saveRDS()` a `save()` je v tom, že funkcia `saveRDS()` umožňuje uložiť iba jeden konkrétny objekt, pričom je možné pri následnom načítaní tohto objektu funkciou `readRDS()` uložiť tento objekt do objektu s iným názvom. Uložený súbor má koncovku `.rds` a je možné ho otvoriť pomocou funkcie `readRDS()`.

```

1 > saveRDS(object = mocniny, file = "mocniny.rds")
2 > noveData <- readRDS(file = "mocniny.rds")
3 > head(mocniny)
4   pismo cislo naDruhu naTretiu
5 A      A      1        1        1
6 B      B      2        4        8
7 C      C      3        9       27
8 D      D      4       16       64
9 E      E      5       25      125
10 F      F      6       36      216
11 > head(noveData)
12   pismo cislo naDruhu naTretiu
13 A      A      1        1        1
14 B      B      2        4        8
15 C      C      3        9       27
16 D      D      4       16       64
17 E      E      5       25      125
18 F      F      6       36      216
19 > all.equal(mocniny, noveData)
20 [1] TRUE

```

7.3.2 Funkcie save() a load()

Ak chceme v prostredí **R** uložiť viacero objektov naraz, môžeme pri tom použiť funkciu `save()`. Pri následnom načítaní týchto objektov funkciou `load()` sa objekty načítajú do pracovného prostredia s názvami, pod akými boli uložené. Uložený súbor má koncovku `.rda` alebo `.RData` a je možné ho otvoriť pomocou funkcie `load()` v **R**.

```

1 save(mocniny, data09, file = "viacObjektov.rda")

```

Aby sme ilustrovali použitie tejto funkcie, vymažeme si po uložení objektov všetko, čo sa nachádza v našom pracovnom prostredí pomocou funkcie `rm(list = ls())`, ktorú sme definovali v Kapitole 2.

```

1 > rm(list = ls())
2 > print(mocniny)
3 Error in print(mocniny) : object 'mocniny' not found
4 > print(data09)
5 Error in print(data09) : object 'data09' not found
6 > load(file = "viacObjektov.rda")
7 > head(mocniny)
8   pismo cislo naDruhu naTretiu
9 A      A      1        1        1
10 B      B      2        4        8

```



```

11 C      C      3      9      27
12 D      D      4     16     64
13 E      E      5     25    125
14 F      F      6     36    216
15 > head(data09)
16 # A tibble: 6 x 2
17       x      y
18   <dbl> <dbl>
19 1  12.4  11.2
20 2  14.3  12.5
21 3  14.5  12.7
22 4  14.9  13.1
23 5  16.1  14.1
24 6  16.9  14.8

```

7.3.3 Funkcie `save.image()` a `load()`

V prípade, ak vykonávame komplexnú analýzu, aby sme nemuseli skript spúšťať nanovo, môžeme si uložiť kompletne všetky aktuálne uložené objekty v celom pracovnom prostredí. Služi na to funkcia `save.image()`. Pri následnom načítaní celého pracovného prostredia využívame funkciu `load()` a všetky objekty sa načítajú do pracovného prostredia s názvami, pod akými boli pôvodne uložené. Uložený súbor má koncovku `.rda` alebo `.RData` a je možné ho otvoriť pomocou funkcie `load()` v **R**.

```

1 > save.image(file = "mojePracovneProstredie.rda")
2 > load(file = "mojePracovneProstredie.rda")

```

7.4 Práca s chýbajúcimi pozorovaniami v R

Mnohé štatistické metódy pri svojom použití predpokladajú, že v analyzovaných údajoch nemáme chýbajúce pozorovania. V praxi sa však v reálnych dátach často stretávame s problémom tzv. chýbajúcich pozorovaní (missing data). Existuje množstvo prístupov, ako riešiť tento problém, napríklad môžeme všetky riadky dátovej matice, v ktorej sa nachádzajú chýbajúce pozorovania vynechať, prípadne môžeme chýbajúce pozorovania nahradiť priemerom za celý stĺpec a podobne. Všetko závisí od charakteru premennej a príčiny chýbajúcej hodnoty, ktorá môže vyplývať napríklad aj z povahy danej premennej. Napríklad ak máme dotazník, v ktorom sa pýtame na príjem, môže sa stať, že respondenti na danú otázku nechcú odpovedať a chýbajúce pozorovania sú dané práve tým. V takomto prípade je riešením vytvoriť nejakú intervalovú premennú, pri ktorej majú respondenti väčšiu tendenciu odpovedať ako v prípade dotazníka, kde majú hodnoty dopisovať.

Pri práci s chýbajúcimi pozorovaniami sú veľmi užitočné funkcie:

- `is.na()`: Funkcia vráti atomický objekt logických hodnôt, pričom v prípade chýbajúcich pozorovaní vráti `TRUE`, inak vráti `FALSE`.
- `complete.cases()`: Funkcia vráti logickú hodnotu `FALSE` pre tie riadky matice alebo objektu typu data frame, ktoré obsahujú aspoň jednu chýbajúcu hodnotu.
- `na.omit()`: Funkcia odstráni všetky riadky matice alebo objektu typu data frame, ktoré obsahujú aspoň jednu chýbajúcu hodnotu.

Mnohé funkcie majú argument `na.rm`, ktorý slúži na ignorovanie chýbajúcich pozorovaní pri výpočte funkcie. Uvažujme napríklad o nasledujúcom objekte typu data frame, ktorý uložíme do objektu s názvom `novyDF`:

```
1 > novyDF <- data.frame(a = c(2, 3, 5, 8),
2 +                       b = c(3, 8, NA, 5),
3 +                       c = c(10, 4, 6, 11))
4 > print(novyDF)
5   a  b  c
6 1 2  3 10
7 2 3  8  4
8 3 5 NA  6
9 4 8  5 11
10 > dim(novyDF)
11 [1] 4 3
```

Vidíme, že data frame uložený v objekte `novyDF` má štyri riadky a tri stĺpce, pričom na prvý pohľad je jasné, že hodnota v treťom riadku a druhom stĺpci je nedostupná (`NA`). Ak by sme chceli napríklad vypočítať aritmetický priemer jeho jednotlivých stĺpcov, použili by sme na to známu funkciu `colMeans()`:

```
1 > colMeans(novyDF)
2   a    b    c
3 4.50  NA 7.75
```

Vidíme, že hodnota `NA` v druhom stĺpci spôsobila, že nie je možné vypočítať aritmetický priemer tohto stĺpca. Pridaním argumentu `na.rm = TRUE` ignorujeme pri výpočte túto hodnotu a aritmetický priemer vypočítame zo všetkých ostatných hodnôt:

```
1 > colMeans(novyDF, na.rm = TRUE)
2   a    b    c
3 4.500000 5.333333 7.750000
4 > mean(novyDF$b, na.rm = TRUE)
5 [1] 5.333333
6 > novyDF$b
7 [1]  3  8 NA  5
8 > (3 + 8 + 5) / 3
9 [1] 5.333333
```

Ak by sme chceli nájsť všetky chýbajúce pozorovania, tak v prípade objektu typu data frame malého rozmeru je to veľmi jednoduché aj vizuálne, avšak pri väčších datasetoch je veľmi vhodné použitie funkcie `is.na()`:

```
1 > is.na(novyDF)
2           a           b           c
3 [1,] FALSE FALSE FALSE
4 [2,] FALSE FALSE FALSE
5 [3,] FALSE  TRUE FALSE
6 [4,] FALSE FALSE FALSE
```

Kombináciou tejto funkcie a odkazovania sa na konkrétne prvky objektu typu data frame môžeme veľmi jednoducho nahradiť všetky hodnoty s hodnotou NA nejakou konkrétnou hodnotou (napríklad nulou), treba však upozorniť na to, že nahradenie chýbajúcich hodnôt nulou môže byť najmä v kontexte vykonávania analýz veľmi nebezpečné. Problematika riešenia problému chýbajúcich pozorovaní je veľmi komplexná a priestor tejto knihy nám neumožňuje sa jej venovať podrobnejšie.

```
1 > novyDF[is.na(novyDF)] <- 0
2 > print(novyDF)
3   a b  c
4 1 2 3 10
5 2 3 8  4
6 3 5 0  6
7 4 8 5 11
```

Vytvoríme si objekt `novyDF` ešte raz, aby znova obsahoval chýbajúce pozorovanie, aby sme mohli na jeho príklade prezentovať aj použitie ďalších funkcií:

```
1 > novyDF <- data.frame(a = c(2, 3, 5, 8),
2 +                       b = c(3, 8, NA, 5),
3 +                       c = c(10, 4, 6, 11))
4 > complete.cases(novyDF)
5 [1]  TRUE  TRUE FALSE  TRUE
```

Funkcia `complete.cases()` vrátila hodnotu `TRUE` pre riadky, ktoré neobsahujú chýbajúce pozorovania. Vidíme, že hodnota `FALSE` je uvedená pri treťom riadku, ktorý obsahuje chýbajúce pozorovanie. Ak by sme chceli celý tento riadok vymazať, môžeme využiť funkciu `na.omit()`:

```
1 > na.omit(novyDF)
2   a b  c
3 1 2 3 10
4 2 3 8  4
5 4 8 5 11
```

Rovnaký efekt dosiahneme aj kombináciou funkcie `complete.cases()` a odkazovania sa na konkrétne prvky objektu typu data frame. V tomto prípade chceme tie riadky, ktoré neobsahujú chýbajúce pozorovania a všetky im odpovedajúce stĺpce:

```
1 > novyDF[complete.cases(novyDF), ]  
2   a b  c  
3 1 2 3 10  
4 2 3 8  4  
5 4 8 5 11
```

Kapitola 8

Práca s dátovými súbormi

S datasetmi sa v **R** stretávame jednak v balíku **datasets**, ktorý je popísaný nižšie, ale **datasets** obsahuje aj množstvo iných balíkov, ktoré do **R** inštalujeme a spúšťame. Tieto balíky predstavujú súbor funkcií, ktoré sú zacielené na konkrétnu oblasť. Práve na týchto datasetoch sa v demonštračných kódoch a príkladoch snažia autori balíku vysvetliť použitie konkrétnej funkcie. Napríklad súčasťou balíka **ggplot2** je viacero datasetov ako napríklad:

- **diamonds**: cena a iné atribúty vyše 50 000 diamantov,
- **economics**: ekonomické časové rady USA na mesačnej báze v rokoch 1967–2015,
- **presidential**: zoznam 11 prezidentov USA spolu s obdobiami, v ktorých boli vo funkcii, a stranou, ktorá ich nominovala od roku 1953 do roku 2017.

K načítaniu týchto dát slúži funkcia **data()**, ktorá nevytvorí žiaden objekt, iba načíta dáta na pozadí.

```

1 > data("presidential", package = "ggplot2")
2 > print(presidential)
3 # A tibble: 11 x 4
4   name      start      end      party
5   <chr>    <date>    <date>    <chr>
6 1 Eisenhower 1953-01-20 1961-01-20 Republican
7 2 Kennedy    1961-01-20 1963-11-22 Democratic
8 3 Johnson    1963-11-22 1969-01-20 Democratic
9 4 Nixon      1969-01-20 1974-08-09 Republican
10 5 Ford       1974-08-09 1977-01-20 Republican
11 6 Carter     1977-01-20 1981-01-20 Democratic
12 7 Reagan     1981-01-20 1989-01-20 Republican
13 8 Bush       1989-01-20 1993-01-20 Republican

```

14	9	Clinton	1993-01-20	2001-01-20	Democratic
15	10	Bush	2001-01-20	2009-01-20	Republican
16	11	Obama	2009-01-20	2017-01-20	Democratic

8.1 Balík datasets a jeho vybrané datasety

Súčasnou štandardnej inštalácie **R** je balík **datasets**, ktorý obsahuje množstvo zaujímavých databáz, ktoré možno automaticky načítať do **R** bez toho, aby sme museli tieto údaje sťahovať a ukladať do pracovného adresára. Ak si chceme vypísať zoznam týchto datasetov aj s ich stručným popisom, môžeme na to použiť funkciu `data()`:

```
1 > data()
```

Zoznam datasetov sa zobrazí priamo do okna skriptu. Vybrané datasety, ktoré sa často používajú pri vyučovaní štatistiky:

- **AirPassengers**: Dataset s mesačnými údajmi o celkovom počte cestujúcich aerolinkami v rokoch 1949 až 1960, ktorým sa dá ilustrovať analýza sezónnosti v časových radoch.
- **mtcars**: Dataset z roku 1974 uverejnený magazínom Motor Trend US obsahujúci 32 pozorovaní áut konkrétnych značiek a 11 ukazovateľov týchto áut (napríklad spotreba, akcelerácia a podobne), ktorým sa často ilustruje použitie vizualizačných metód, prípadne metódy viacnásobnej regresnej analýzy.
- **ToothGrowth**: Dataset obsahujúci 60 pozorovaní troch premenenných pri testovaní efektu vitamínu C na rast zubov guinejských prasiat. Výsledky tohto experimentu vykonaného v roku 1947 sú vhodné na ilustráciu porovnávania dvoch malých nezávislých normálne rozdelených vzoriek a praktickou aplikáciou Studentovho t-testu, nakoľko jedna z premenných je kategorická premenná nadobúdajúca hodnoty VC (vitamin C), pokiaľ bola prasatám podávaná kyselina askorbová, a hodnoty OJ (orange juice), pokiaľ bol prasatám podávaný pomarančový džús.
- **iris**: Dataset z roku 1935, ktorý obsahuje 150 rastlín (kosatcov) troch rôznych druhov (*Iris Setosa*, *Iris Versicolour*, *Iris Virginica*), pričom každý druh je zhodne zastúpený 50 objektami a je charakterizovaný pomocou štyroch kvantitatívnych premenných: dĺžka a šírka okvetného lístku a dĺžka a šírka kališného lístku. Tento známy dataset sa často používa pri klasifikačných analýzach ako napríklad zhluková analýza, diskriminačná analýza a podobne.
- **PlantGrowth**: Dataset obsahujúci 30 výsledkov o sušenej hmotnosti rastlín rozdelených do troch skupín na základe toho, či išlo o kontrolnú vzorku alebo o aplikáciu jednej z dvoch metód starostlivosti o rastliny, sa často používa k vysvetleniu takzvanej analýzy rozptylu.

- **USArrests**: Dataset z roku 1973 uvádza počet páchateľov trestných činov napadnutia, vraždy a znásilnenia pripadajúci na 100 000 obyvateľov všetkých 50 štátov USA. Okrem toho je v datasete uvedené aj percento populácie, ktoré žije v mestských oblastiach daných štátov. Dataset je vhodný najmä na takzvanú priestorovú analýzu (spatial analysis).

8.2 Praktické spracovanie dát z datasetu mtcars

Dokumentácia k datasetom funguje rovnako ako dokumentácia k funkciám. Po načítaní konkrétneho datasetu je možné zobraziť jeho dokumentáciu prostredníctvom funkcie `help()` alebo symbolu `?`:

```
1 > ?mtcars
2 > help(mtcars)
```

Ak by sme si vypísali túto dokumentáciu, zistili by sme, že daný dataset obsahuje údaje publikované v roku 1974 magazínom Motor Trend US. Riadkami (pozorovaniami) tohto datasetu sú jednotlivé modely áut a stĺpcami tohto datasetu sú jednotlivé premenné. Zoznam premenných je nasledovný:

- **mpg**: Miles/(US) gallon (spotreba),
- **cyl**: Number of cylinders (počet valcov),
- **disp**: Displacement (zdvihový objem motora),
- **hp**: Gross horsepower (hrubý výkon, v počte koní),
- **drat**: Rear axle ratio (prevodový pomer zadnej nápravy),
- **wt**: Weight (hmotnosť),
- **qsec**: 1/4 mile time (čas, za ktorý auto prejde 1/4 míle, kvázi zrýchlenie),
- **vs**: Engine (0 = V-shaped, 1 = straight) (binárna premenná vyjadrujúca tvar motora, 0 znamená tvar písmena V a 1 znamená tvar priamky),
- **am**: Transmission (0 = automatic, 1 = manual) (binárna premenná predstavujúca typ prevodovky, 0 znamená automatickú prevodovku a 1 znamená manuálnu prevodovku),
- **gear**: Number of forward gears (počet prevodových stupňov),
- **carb**: Number of carburetors (počet karburátorov).

Kto má záujem o podrobnejšiu diskusiu týchto premenných, odporúčame nasledujúci zdroj ([odkaz tu](#)).

Príklad 1

Načítajte si údaje z datasetu `mtcars` do objektu `auta` tak, aby menovky riadkov predstavovali názvy jednotlivých áut, a určte, koľko pozorovaní a koľko premenných má analyzovaný dataset.

Riešenie 1

Objekt `auta` vytvoríme priradením datasetu `mtcars` do daného objektu, menovky riadkov sú v prípade vstavaného datasetu nastavené automaticky. Na určenie počtu pozorovaní môžeme použiť funkciu `nrow()` a na určenie počtu premenných môžeme použiť funkciu `ncol()`, prípadne súhrnne môžeme použiť funkciu `dim()`:

```
1 > auta <- mtcars
2 > row.names(auta)
3 [1] "Mazda RX4"           "Mazda RX4 Wag"
4 [3] "Datsun 710"          "Hornet 4 Drive"
5 [5] "Hornet Sportabout"   "Valiant"
6 [7] "Duster 360"          "Merc 240D"
7 [9] "Merc 230"            "Merc 280"
8 [11] "Merc 280C"           "Merc 450SE"
9 [13] "Merc 450SL"          "Merc 450SLC"
10 [15] "Cadillac Fleetwood"  "Lincoln Continental"
11 [17] "Chrysler Imperial"  "Fiat 128"
12 [19] "Honda Civic"         "Toyota Corolla"
13 [21] "Toyota Corona"      "Dodge Challenger"
14 [23] "AMC Javelin"         "Camaro Z28"
15 [25] "Pontiac Firebird"    "Fiat X1-9"
16 [27] "Porsche 914-2"       "Lotus Europa"
17 [29] "Ford Pantera L"      "Ferrari Dino"
18 [31] "Maserati Bora"       "Volvo 142E"
19 > ncol(auta)
20 [1] 11
21 > nrow(auta)
22 [1] 32
23 > dim(auta)
24 [1] 32 11
```

Vidíme, že v datasete máme 32 pozorovaní (áut) a 11 premenných.

Príklad 2

Vypočítajte súhrnné popisné štatistiky objektu typu data frame uloženého v objekte `auta`. Majú tieto popisné štatistiky zmysel pre všetky premenné?

Riešenie 2

```

1 > summary(auta)
2      mpg           cyl          disp           hp
3  Min.    :10.40   Min.    : 4.000   Min.    : 71.1   Min.    : 52.0
4  1st Qu.:15.43   1st Qu.: 4.000   1st Qu.:120.8   1st Qu.: 96.5
5  Median :19.20   Median : 6.000   Median :196.3   Median :123.0
6  Mean    :20.09   Mean    : 6.188   Mean    :230.7   Mean    :146.7
7  3rd Qu.:22.80   3rd Qu.: 8.000   3rd Qu.:326.0   3rd Qu.:180.0
8  Max.    :33.90   Max.    : 8.000   Max.    :472.0   Max.    :335.0
9
10      drat          wt          qsec           vs
11  Min.    :2.760   Min.    :1.513   Min.    :14.50   Min.    :0.0000
12  1st Qu.:3.080   1st Qu.:2.581   1st Qu.:16.89   1st Qu.:0.0000
13  Median :3.695   Median :3.325   Median :17.71   Median :0.0000
14  Mean    :3.597   Mean    :3.217   Mean    :17.85   Mean    :0.4375
15  3rd Qu.:3.920   3rd Qu.:3.610   3rd Qu.:18.90   3rd Qu.:1.0000
16  Max.    :4.930   Max.    :5.424   Max.    :22.90   Max.    :1.0000
17
18      am          gear          carb
19  Min.    :0.0000   Min.    :3.000   Min.    :1.000
20  1st Qu.:0.0000   1st Qu.:3.000   1st Qu.:2.000
21  Median :0.0000   Median :4.000   Median :2.000
22  Mean    :0.4062   Mean    :3.688   Mean    :2.812
23  3rd Qu.:1.0000   3rd Qu.:4.000   3rd Qu.:4.000
24  Max.    :1.0000   Max.    :5.000   Max.    :8.000

```

Všimnime si, že v stĺpcoch `vs` a `am` sa nachádzajú binárne premenné vyjadrujúce tvar motora, respektíve typ prevodovky. Tým, že sú „zakódované“ číselnými hodnotami 0 (tvar motora V, respektíve automatická prevodovka) a 1 (tvar motora priamka, respektíve manuálna prevodovka), **R** ich berie ako čísla a počíta s nimi klasicky ako s číslami. Je však jasné, že nemá zmysel pri týchto kategorických binárnych premenných počítať číselné charakteristiky, ktoré využívame pri číselných spojitých premenných ako napríklad aritmetický priemer. Priemerná hodnota premennej v stĺpci `am`, ktorá sa vzťahuje na typ prevodovky na úrovni 0,4062 vypovedá jedine o tom, že v súbore je viac áut s automatickou prevodovkou ako s manuálnou prevodovkou. Z toho vyplýva, že musíme **R** povedať, že daná premenná nie je číselná, aj keď jej hodnoty sú číselné, ale ide o nominálnu kategorickú premennú. To môžeme urobiť prostredníctvom funkcie `as.factor()`:

```

1 > auta$vs <- as.factor(auta$vs)
2 > auta$am <- as.factor(auta$am)
3 > summary(auta)
4      mpg           cyl          disp           hp
5  Min.    :10.40   Min.    : 4.000   Min.    : 71.1   Min.    : 52.0
6  1st Qu.:15.43   1st Qu.: 4.000   1st Qu.:120.8   1st Qu.: 96.5
7  Median :19.20   Median : 6.000   Median :196.3   Median :123.0
8  Mean    :20.09   Mean    : 6.188   Mean    :230.7   Mean    :146.7
9  3rd Qu.:22.80   3rd Qu.: 8.000   3rd Qu.:326.0   3rd Qu.:180.0

```

```

10 Max.      :33.90    Max.      :8.000    Max.      :472.0    Max.      :335.0
11      drat          wt          qsec          vs          am
12 Min.      :2.760    Min.      :1.513    Min.      :14.50    0:18      0:19
13 1st Qu.:3.080    1st Qu.:2.581    1st Qu.:16.89    1:14      1:13
14 Median :3.695    Median :3.325    Median :17.71
15 Mean     :3.597    Mean     :3.217    Mean     :17.85
16 3rd Qu.:3.920    3rd Qu.:3.610    3rd Qu.:18.90
17 Max.     :4.930    Max.     :5.424    Max.     :22.90
18      gear          carb
19 Min.      :3.000    Min.      :1.000
20 1st Qu.:3.000    1st Qu.:2.000
21 Median :4.000    Median :2.000
22 Mean     :3.688    Mean     :2.812
23 3rd Qu.:4.000    3rd Qu.:4.000
24 Max.     :5.000    Max.     :8.000

```

Vidíme, že pokiaľ ide o kategorickú premennú, funkcia `summary()` namiesto popisných štatistík vypočíta absolútne početnosti (v ČJ četnosti) jednotlivých kategórií. Vidíme napríklad, že v súbore máme 18 áut s motorom v tvare písmena V a 14 áut s motorom v tvare priamky. Ďalej vidíme, že v súbore naozaj máme viac áut s automatickou prevodovkou, konkrétne 19 áut s manuálnou prevodovkou je 13.

Príklad 3

Vytvorte tabuľku početností pre premenné uložené v stĺpcoch `cyl` a `gear` predstavujúce počet valcov a počet prevodových stupňov.

Riešenie 3

Tabuľku početností vieme vytvoriť pomocou funkcie `table()`:

```

1 > table(auta$cyl)
2
3  4   6   8
4 11   7  14
5 > table(auta$gear)
6
7  3   4   5
8 15  12   5

```

Ako môžeme vidieť, autá majú buď štyri, šesť, alebo osem valcov, pričom dominujú osemvalcové autá. Čo sa týka počtu prevodových stupňov, môže ich byť tri, štyri alebo päť, pričom v tomto prípade dominujú autá, ktoré majú tri prevodové stupne.

Príklad 4

V podrobnejšej dokumentácii k danému datasetu je uvedené, že manuálne prevodovky majú buď štyri, alebo päť prevodových stupňov a že automatické prevodovky majú buď tri, alebo štyri prevodové stupne. Overme to prostredníctvom kontingenčnej tabuľky.

Riešenie 4

Kontingenčnú tabuľku vieme vytvoriť taktiež pomocou funkcie `table()`:

```
1 > table(auta$am, auta$gear)
2
3      3  4  5
4  0 15  4  0
5  1  0  8  5
```

V riadkoch tejto matice je uvedený typ prevodovky (0: automatická, 1: manuálna) a v stĺpcoch tejto matice je uvedený počet prevodových stupňov.

Vidíme, že pätnásť áut, ktoré majú automatickú prevodovku, má tri prevodové stupne. Štyri autá, ktoré majú automatickú prevodovku, majú štyri prevodové stupne. V súbore sa nenachádza auto, ktoré by malo automatickú prevodovku a päť prevodových stupňov.

Vidíme, že osem áut, ktoré majú manuálnu prevodovku, má štyri prevodové stupne. Päť áut, ktoré majú manuálnu prevodovku, má päť prevodových stupňov. V súbore sa nenachádza auto, ktoré by malo manuálnu prevodovku a tri prevodové stupne.

Na základe kontingenčnej tabuľky sme ukázali, že manuálne prevodovky majú buď štyri, alebo päť prevodových stupňov a že automatické prevodovky majú buď tri, alebo štyri prevodové stupne.

Príklad 5

Nájdite všetky informácie o automobile `Lincoln Continental`.

Riešenie 5

```
1 > rownames(auta)
2 [1] "Mazda RX4"           "Mazda RX4 Wag"
3 [3] "Datsun 710"          "Hornet 4 Drive"
4 [5] "Hornet Sportabout"   "Valiant"
5 [7] "Duster 360"          "Merc 240D"
6 [9] "Merc 230"            "Merc 280"
7 [11] "Merc 280C"           "Merc 450SE"
8 [13] "Merc 450SL"          "Merc 450SLC"
```

```

9 [15] "Cadillac Fleetwood" "Lincoln Continental"
10 [17] "Chrysler Imperial" "Fiat 128"
11 [19] "Honda Civic" "Toyota Corolla"
12 [21] "Toyota Corona" "Dodge Challenger"
13 [23] "AMC Javelin" "Camaro Z28"
14 [25] "Pontiac Firebird" "Fiat X1-9"
15 [27] "Porsche 914-2" "Lotus Europa"
16 [29] "Ford Pantera L" "Ferrari Dino"
17 [31] "Maserati Bora" "Volvo 142E"
18 > rownames(auta) == "Lincoln Continental"
19 [1] FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE
20 [11] FALSE FALSE FALSE FALSE FALSE TRUE FALSE FALSE FALSE FALSE
21 [21] FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE FALSE
22 [31] FALSE FALSE
23 > which(rownames(auta) == "Lincoln Continental")
24 [1] 16
25 > auta[which(rownames(auta) == "Lincoln Continental"), ]
26           mpg cyl disp  hp drat   wt  qsec vs am gear
27 Lincoln Continental 10.4   8  460 215   3 5.424 17.82  0  0    3
28           carb
29 Lincoln Continental    4

```

Najprv sme si funkciou `rownames()` iba vypísali mená všetkých áut, potom sme prostredníctvom logického operátora porovnania (`==`) chceli vedieť, ktoré autá sa volajú *Lincoln Continental*. Iba jedno z nich sa tak volá, preto je každá hodnota `FALSE` s výnimkou jednej, ktorá je `TRUE`. Ak chceme poznať poradie tejto hodnoty, môžeme použiť funkciu `which()`, ktorá z logického vektora vyberie pozície, na ktorých sa nachádzajú hodnoty `TRUE`. V našom prípade ide o pozíciu číslo 16, čo znamená, že nami hľadaný automobil *Lincoln Continental* sa nachádza v 16. riadku. Následne už iba vypíšeme tento riadok a všetky stĺpce.

Riešiť tento príklad je možné aj jednoduchším spôsobom pomocou úvodzoviek, nakoľko máme zadefinované menovky riadkov:

```

1 > auta["Lincoln Continental", ]
2           mpg cyl disp  hp drat   wt  qsec vs am gear
3 Lincoln Continental 10.4   8  460 215   3 5.424 17.82  0  0    3
4           carb
5 Lincoln Continental    4

```

Príklad 6

Zistite, ktoré tri autá majú najnižšiu spotrebu.

Riešenie 6

Najnižšiu spotrebu majú autá, ktoré sú schopné prejsť najviac míľ na jeden galón paliva. Spotrebu máme uloženú v stĺpci mpg:

```
1 > auta$mpg
2 [1] 21.0 21.0 22.8 21.4 18.7 18.1 14.3 24.4 22.8 19.2 17.8 16.4
3 [13] 17.3 15.2 10.4 10.4 14.7 32.4 30.4 33.9 21.5 15.5 15.2 13.3
4 [25] 19.2 27.3 26.0 30.4 15.8 19.7 15.0 21.4
5 > order(auta$mpg, decreasing = TRUE)
6 [1] 20 18 19 28 26 27 8 3 9 21 4 32 1 2 30 10 25 5 6 11
7 [21] 13 12 29 22 14 23 31 17 7 24 15 16
8 > order(auta$mpg, decreasing = TRUE)[1:3]
9 [1] 20 18 19
```

Najprv si iba vypíšeme hodnoty spotreby jednotlivých automobilov. Hľadáme, ktoré tri z nich majú túto hodnotu najvyššiu. Môžeme využiť funkciu `order()` s argumentom `decreasing = TRUE` (zostupne), ktorá hodnoty usporiada od najväčšej po najmenšiu a určí poradie týchto hodnôt. Môžeme vidieť, že najnižšiu spotrebu, a teda najviac míľ na jeden galón paliva, je schopné prejsť auto, ktoré je 20. v poradí. Druhú najnižšiu spotrebu má auto, ktoré je 18. v poradí, a tretiu najnižšiu spotrebu má auto, ktoré je 19. v poradí. Môžeme si vypísať ich mená a hodnoty jednotlivých ukazovateľov.

```
1 > rownames(auta)[order(auta$mpg, decreasing = TRUE)[1:3]]
2 [1] "Toyota Corolla" "Fiat 128" "Honda Civic"
3 > which(colnames(auta) == "mpg")
4 [1] 1
5 > auta[order(auta$mpg, decreasing = TRUE)[1:3], 1]
6 [1] 33.9 32.4 30.4
7 > auta[order(auta$mpg, decreasing = TRUE)[1:3], ]
8      mpg cyl disp hp drat   wt  qsec vs am gear carb
9 Toyota Corolla 33.9  4 71.1 65 4.22 1.835 19.90 1 1  4    1
10 Fiat 128      32.4  4 78.7 66 4.08 2.200 19.47 1 1  4    1
11 Honda Civic   30.4  4 75.7 52 4.93 1.615 18.52 1 1  4    2
```

Vidíme, že ide o automobily Toyota Corolla, Fiat 128 a Honda Civic. Počet prejdejších míľ na jeden galón paliva je uložený v stĺpci s názvom mpg, ktorý je prvý v poradí, preto vidíme, že najviac míľ na galón paliva dokáže prejsť Toyota Corolla (33,9), za ňou nasleduje Fiat 128 (32,4) a tretia v poradí je Honda Civic (30,4).

Príklad 7

Zistite, ktoré tri autá majú najväčšie zrýchlenie.

Riešenie 7

Najväčšie zrýchlenie majú autá, ktoré dokážu prejsť štvrtinu míle za najkratší čas. Hľadáme preto tri najmenšie hodnoty v stĺpci `qsec`:

```
1 > auta$qsec
2 [1] 16.46 17.02 18.61 19.44 17.02 20.22 15.84 20.00 22.90 18.30
3 [11] 18.90 17.40 17.60 18.00 17.98 17.82 17.42 19.47 18.52 19.90
4 [21] 20.01 16.87 17.30 15.41 17.05 18.90 16.70 16.90 14.50 15.50
5 [31] 14.60 18.60
6 > order(auta$qsec, decreasing = FALSE)
7 [1] 29 31 24 30 7 1 27 22 28 2 5 25 23 12 17 13 16 15 14 10
8 [21] 19 32 3 11 26 4 18 20 8 21 6 9
9 > order(auta$qsec, decreasing = FALSE)[1:3]
10 [1] 29 31 24
11 >
12 > rownames(auta)[order(auta$qsec, decreasing = FALSE)[1:3]]
13 [1] "Ford Pantera L" "Maserati Bora" "Camaro Z28"
14 > which(colnames(auta) == "qsec")
15 [1] 7
16 > auta[order(auta$qsec, decreasing = FALSE)[1:3], 7]
17 [1] 14.50 14.60 15.41
18 > auta[order(auta$qsec, decreasing = FALSE)[1:3], ]
19      mpg cyl disp  hp drat   wt  qsec vs am gear carb
20 Ford Pantera L 15.8   8  351 264 4.22 3.17 14.50 0  1   5    4
21 Maserati Bora  15.0   8  301 335 3.54 3.57 14.60 0  1   5    8
22 Camaro Z28     13.3   8  350 245 3.73 3.84 15.41 0  0   3    4
```

Môžeme využiť priamo funkciu `order()` na usporiadanie hodnôt v tabuľke podľa konkrétneho stĺpca:

```
1 > auta[order(auta$qsec), ]
2      mpg cyl disp  hp drat   wt  qsec vs am gear carb
3 Ford Pantera L 15.8   8 351.0 264 4.22 3.170 14.50 0  1   5
4 Maserati Bora  15.0   8 301.0 335 3.54 3.570 14.60 0  1   5
5 Camaro Z28     13.3   8 350.0 245 3.73 3.840 15.41 0  0   3
6 Ferrari Dino   19.7   6 145.0 175 3.62 2.770 15.50 0  1   5
7 Duster 360     14.3   8 360.0 245 3.21 3.570 15.84 0  0   3
8 Mazda RX4     21.0   6 160.0 110 3.90 2.620 16.46 0  1   4
9 Porsche 914-2  26.0   4 120.3  91 4.43 2.140 16.70 0  1   5
10 Dodge Challenger 15.5   8 318.0 150 2.76 3.520 16.87 0  0   3
11 Lotus Europa   30.4   4  95.1 113 3.77 1.513 16.90 1  1   5
12 Mazda RX4 Wag  21.0   6 160.0 110 3.90 2.875 17.02 0  1   4
13 Hornet Sportabout 18.7   8 360.0 175 3.15 3.440 17.02 0  0   3
14 Pontiac Firebird 19.2   8 400.0 175 3.08 3.845 17.05 0  0   3
15 AMC Javelin    15.2   8 304.0 150 3.15 3.435 17.30 0  0   3
16 Merc 450SE     16.4   8 275.8 180 3.07 4.070 17.40 0  0   3
```

17	Chrysler Imperial	14.7	8	440.0	230	3.23	5.345	17.42	0	0	3
18	Merc 450SL	17.3	8	275.8	180	3.07	3.730	17.60	0	0	3
19	Lincoln Continental	10.4	8	460.0	215	3.00	5.424	17.82	0	0	3
20	Cadillac Fleetwood	10.4	8	472.0	205	2.93	5.250	17.98	0	0	3
21	Merc 450SLC	15.2	8	275.8	180	3.07	3.780	18.00	0	0	3
22	Merc 280	19.2	6	167.6	123	3.92	3.440	18.30	1	0	4
23	Honda Civic	30.4	4	75.7	52	4.93	1.615	18.52	1	1	4
24	Volvo 142E	21.4	4	121.0	109	4.11	2.780	18.60	1	1	4
25	Datsun 710	22.8	4	108.0	93	3.85	2.320	18.61	1	1	4
26	Merc 280C	17.8	6	167.6	123	3.92	3.440	18.90	1	0	4
27	Fiat X1-9	27.3	4	79.0	66	4.08	1.935	18.90	1	1	4
28	Hornet 4 Drive	21.4	6	258.0	110	3.08	3.215	19.44	1	0	3
29	Fiat 128	32.4	4	78.7	66	4.08	2.200	19.47	1	1	4
30	Toyota Corolla	33.9	4	71.1	65	4.22	1.835	19.90	1	1	4
31	Merc 240D	24.4	4	146.7	62	3.69	3.190	20.00	1	0	4
32	Toyota Corona	21.5	4	120.1	97	3.70	2.465	20.01	1	0	3
33	Valiant	18.1	6	225.0	105	2.76	3.460	20.22	1	0	3
34	Merc 230	22.8	4	140.8	95	3.92	3.150	22.90	1	0	4
35		carb									
36	Ford Pantera L	4									
37	Maserati Bora	8									
38	Camaro Z28	4									
39	Ferrari Dino	6									
40	Duster 360	4									
41	Mazda RX4	4									
42	Porsche 914-2	2									
43	Dodge Challenger	2									
44	Lotus Europa	2									
45	Mazda RX4 Wag	4									
46	Hornet Sportabout	2									
47	Pontiac Firebird	2									
48	AMC Javelin	2									
49	Merc 450SE	3									
50	Chrysler Imperial	4									
51	Merc 450SL	3									
52	Lincoln Continental	4									
53	Cadillac Fleetwood	4									
54	Merc 450SLC	3									
55	Merc 280	4									
56	Honda Civic	2									
57	Volvo 142E	2									
58	Datsun 710	1									
59	Merc 280C	4									
60	Fiat X1-9	1									
61	Hornet 4 Drive	1									
62	Fiat 128	1									

```

63 Toyota Corolla      1
64 Merc 240D          2
65 Toyota Corona       1
66 Valiant             1
67 Merc 230            2

```

Ak by sme chceli poznať konkrétne usporiadané hodnoty, môžeme na to použiť funkciu `sort()`:

```

1 > sort(auta$qsec)
2 [1] 14.50 14.60 15.41 15.50 15.84 16.46 16.70 16.87 16.90 17.02
3 [11] 17.02 17.05 17.30 17.40 17.42 17.60 17.82 17.98 18.00 18.30
4 [21] 18.52 18.60 18.61 18.90 18.90 19.44 19.47 19.90 20.00 20.01
5 [31] 20.22 22.90

```

Príklad 8

Zistite základné informácie o zdvihovom objeme motora a vytvorte v datasete nový stĺpec, v ktorom rozdelíte hodnoty zdvihového objemu motora do štyroch skupín:

- do 100 kubických palcov vrátane,
- nad 100 do 200 kubických palcov vrátane,
- nad 200 do 300 kubických palcov vrátane,
- nad 300 kubických palcov.

Následne vytvorte tabuľku početností novovytvoreného ukazovateľa a kontingenčnú tabuľku novovytvoreného ukazovateľa a tvaru motora. Čo by sa dalo z tejto tabuľky predpokladať?

Riešenie 8

```

1 > auta$disp
2 [1] 160.0 160.0 108.0 258.0 360.0 225.0 360.0 146.7 140.8 167.6
3 [11] 167.6 275.8 275.8 275.8 472.0 460.0 440.0 78.7 75.7 71.1
4 [21] 120.1 318.0 304.0 350.0 400.0 79.0 120.3 95.1 351.0 145.0
5 [31] 301.0 121.0
6 > min(auta$disp)
7 [1] 71.1
8 > max(auta$disp)
9 [1] 472

```


Zdvihový objem motora sa nachádza v stĺpci `disp`. Vidíme, že rozsah hodnôt tohto ukazovateľa je od 71,1 kubických palcov po 472 kubických palcov. Číselne spojitú premennú je možné rozdeliť do intervalov prostredníctvom funkcie `cut()`. Jedným z argumentov tejto funkcie je `breaks`, ktorým môžeme priamo nadefinovať hranice pomocou vektora s hodnotami dolných a horných hraníc intervalov. Ak za argument `breaks` použijeme celé číslo n , vytvorí sa n ekvidištančných intervalov, pričom dolná hranica prvého je rovná približne minimu a horná hranica posledného je rovná približne maximu. Môžeme si to vyskúšať napríklad pri vytvorení piatich intervalov rovnakej šírky:

```
1 > cut(auta$disp, breaks = 5)
2 [1] (151,231] (151,231] (70.7,151] (231,312] (312,392]
3 [6] (151,231] (312,392] (70.7,151] (70.7,151] (151,231]
4 [11] (151,231] (231,312] (231,312] (231,312] (392,472]
5 [16] (392,472] (392,472] (70.7,151] (70.7,151] (70.7,151]
6 [21] (70.7,151] (312,392] (231,312] (312,392] (392,472]
7 [26] (70.7,151] (70.7,151] (70.7,151] (312,392] (70.7,151]
8 [31] (231,312] (70.7,151]
9 Levels: (70.7,151] (151,231] (231,312] (312,392] (392,472]
10 > table(cut(auta$disp, breaks = 5))
11
12 (70.7,151] (151,231] (231,312] (312,392] (392,472]
13      12          5          6          5          4
```

Ak si chceme vytvoriť hranice sami, využívame argument `breaks` nasledovne:

```
1 > cut(
2 +   auta$disp,
3 +   breaks = c(0, 100, 200, 300, 1000),
4 +   labels = c("do 100", "100-200", "200-300", "nad 300")
5 + )
6 [1] 100-200 100-200 100-200 200-300 nad 300 200-300 nad 300
7 [8] 100-200 100-200 100-200 100-200 200-300 200-300 200-300
8 [15] nad 300 nad 300 nad 300 do 100 do 100 do 100 100-200
9 [22] nad 300 nad 300 nad 300 nad 300 do 100 100-200 do 100
10 [29] nad 300 100-200 nad 300 100-200
11 Levels: do 100 100-200 200-300 nad 300
12 >
13 > auta$disp_group <- cut(
14 +   auta$disp,
15 +   breaks = c(0, 100, 200, 300, 1000),
16 +   labels = c("do 100", "100-200", "200-300", "nad 300")
17 + )
18 >
19 > head(cbind.data.frame(auta$disp, auta$disp_group))
20   auta$disp auta$disp_group
21 1        160      100-200
```

22	2	160	100-200
23	3	108	100-200
24	4	258	200-300
25	5	360	nad 300
26	6	225	200-300

Pri tomto vytváraní vlastných intervalov je potrebné upozorniť na to, že ak by sme zobrali za dolnú hranicu minimálnu hodnotu súboru, hodnoty rovné minimu by neboli priradené do žiadneho z intervalov, nakoľko funkcia `cut()` vytvára intervaly typu $(a, b >$. Je to tak z dôvodu, že primárne nastavenie funkcie `cut()` je `include.lowest = FALSE`, čo znamená, že zľava je tento interval otvorený. Preto pri definovaní vlastných hraníc odporúčame radšej ich definovanie od nižšej hodnoty, ako je najnižšia hodnota v súbore po vyššiu hodnotu, ako je najvyššia hodnota v súbore.

Na vytvorenie tabuľky početností a kontingenčnej tabuľky môžeme použiť funkciu `table()`:

```

1 > table(auta$disp_group)
2
3 do 100 100-200 200-300 nad 300
4      5      11      5      11
5 > table(auta$disp_group, auta$vs)
6
7      0  1
8 do 100  0  5
9 100-200  4  7
10 200-300  3  2
11 nad 300 11  0

```

Na prvý pohľad z kontingenčnej tabuľky vidíme, že autá s vyšším zdvihovým objemom motora majú väčšinou motor tvaru písmena V a autá s nižším zdvihovým objemom motora majú väčšinou motor tvaru priamky.

Príklad 9

Nájdite tie automobily, ktorých hmotnosť je väčšia ako aritmetický priemer hmotností všetkých automobilov.

Riešenie 9

```

1 > priemernaHmotnost <- mean(auta$wt)
2 > print(priemernaHmotnost)
3 [1] 3.21725
4 > which(auta$wt > priemernaHmotnost)
5 [1] 5 6 7 10 11 12 13 14 15 16 17 22 23 24 25 31

```

```

6 >
7 > rownames(auta)[which(auta$wt > priemernaHmotnost)]
8 [1] "Hornet Sportabout" "Valiant"
9 [3] "Duster 360" "Merc 280"
10 [5] "Merc 280C" "Merc 450SE"
11 [7] "Merc 450SL" "Merc 450SLC"
12 [9] "Cadillac Fleetwood" "Lincoln Continental"
13 [11] "Chrysler Imperial" "Dodge Challenger"
14 [13] "AMC Javelin" "Camaro Z28"
15 [15] "Pontiac Firebird" "Maserati Bora"
16 > auta[which(auta$wt > priemernaHmotnost),]
17      mpg   cyl  disp    hp  drat    wt   qsec    vs  am
18 Hornet Sportabout  18.7    8 360.0  175  3.15  3.440  17.02    0   0
19 Valiant            18.1    6 225.0  105  2.76  3.460  20.22    1   0
20 Duster 360        14.3    8 360.0  245  3.21  3.570  15.84    0   0
21 Merc 280          19.2    6 167.6  123  3.92  3.440  18.30    1   0
22 Merc 280C         17.8    6 167.6  123  3.92  3.440  18.90    1   0
23 Merc 450SE        16.4    8 275.8  180  3.07  4.070  17.40    0   0
24 Merc 450SL        17.3    8 275.8  180  3.07  3.730  17.60    0   0
25 Merc 450SLC       15.2    8 275.8  180  3.07  3.780  18.00    0   0
26 Cadillac Fleetwood 10.4    8 472.0  205  2.93  5.250  17.98    0   0
27 Lincoln Continental 10.4    8 460.0  215  3.00  5.424  17.82    0   0
28 Chrysler Imperial 14.7    8 440.0  230  3.23  5.345  17.42    0   0
29 Dodge Challenger  15.5    8 318.0  150  2.76  3.520  16.87    0   0
30 AMC Javelin       15.2    8 304.0  150  3.15  3.435  17.30    0   0
31 Camaro Z28        13.3    8 350.0  245  3.73  3.840  15.41    0   0
32 Pontiac Firebird   19.2    8 400.0  175  3.08  3.845  17.05    0   0
33 Maserati Bora      15.0    8 301.0  335  3.54  3.570  14.60    0   1
34      gear carb disp_group
35 Hornet Sportabout    3     2      nad 300
36 Valiant              3     1    200-300
37 Duster 360           3     4      nad 300
38 Merc 280             4     4    100-200
39 Merc 280C            4     4    100-200
40 Merc 450SE           3     3    200-300
41 Merc 450SL           3     3    200-300
42 Merc 450SLC          3     3    200-300
43 Cadillac Fleetwood   3     4      nad 300
44 Lincoln Continental   3     4      nad 300
45 Chrysler Imperial    3     4      nad 300
46 Dodge Challenger     3     2      nad 300
47 AMC Javelin          3     2      nad 300
48 Camaro Z28           3     4      nad 300
49 Pontiac Firebird      3     2      nad 300
50 Maserati Bora         5     8      nad 300
51 >

```

```

52 > rownames(auta[auta$wt > priemernaHmotnost,])
53 [1] "Hornet Sportabout" "Valiant"
54 [3] "Duster 360" "Merc 280"
55 [5] "Merc 280C" "Merc 450SE"
56 [7] "Merc 450SL" "Merc 450SLC"
57 [9] "Cadillac Fleetwood" "Lincoln Continental"
58 [11] "Chrysler Imperial" "Dodge Challenger"
59 [13] "AMC Javelin" "Camaro Z28"
60 [15] "Pontiac Firebird" "Maserati Bora"
61 > auta[auta$wt > priemernaHmotnost,]
62 mpg cyl disp hp drat wt qsec vs am
63 Hornet Sportabout 18.7 8 360.0 175 3.15 3.440 17.02 0 0
64 Valiant 18.1 6 225.0 105 2.76 3.460 20.22 1 0
65 Duster 360 14.3 8 360.0 245 3.21 3.570 15.84 0 0
66 Merc 280 19.2 6 167.6 123 3.92 3.440 18.30 1 0
67 Merc 280C 17.8 6 167.6 123 3.92 3.440 18.90 1 0
68 Merc 450SE 16.4 8 275.8 180 3.07 4.070 17.40 0 0
69 Merc 450SL 17.3 8 275.8 180 3.07 3.730 17.60 0 0
70 Merc 450SLC 15.2 8 275.8 180 3.07 3.780 18.00 0 0
71 Cadillac Fleetwood 10.4 8 472.0 205 2.93 5.250 17.98 0 0
72 Lincoln Continental 10.4 8 460.0 215 3.00 5.424 17.82 0 0
73 Chrysler Imperial 14.7 8 440.0 230 3.23 5.345 17.42 0 0
74 Dodge Challenger 15.5 8 318.0 150 2.76 3.520 16.87 0 0
75 AMC Javelin 15.2 8 304.0 150 3.15 3.435 17.30 0 0
76 Camaro Z28 13.3 8 350.0 245 3.73 3.840 15.41 0 0
77 Pontiac Firebird 19.2 8 400.0 175 3.08 3.845 17.05 0 0
78 Maserati Bora 15.0 8 301.0 335 3.54 3.570 14.60 0 1
79 gear carb disp_group
80 Hornet Sportabout 3 2 nad 300
81 Valiant 3 1 200-300
82 Duster 360 3 4 nad 300
83 Merc 280 4 4 100-200
84 Merc 280C 4 4 100-200
85 Merc 450SE 3 3 200-300
86 Merc 450SL 3 3 200-300
87 Merc 450SLC 3 3 200-300
88 Cadillac Fleetwood 3 4 nad 300
89 Lincoln Continental 3 4 nad 300
90 Chrysler Imperial 3 4 nad 300
91 Dodge Challenger 3 2 nad 300
92 AMC Javelin 3 2 nad 300
93 Camaro Z28 3 4 nad 300
94 Pontiac Firebird 3 2 nad 300
95 Maserati Bora 5 8 nad 300

```

Príklad 10

Nájdite tie automobily, ktorých hrubý výkon je viac ako 200 koní a ktoré majú päťstupňovú prevodovku.

Riešenie 10

```
1 > rownames(auta)[auta$hp > 200]
2 [1] "Duster 360" "Cadillac Fleetwood"
3 [3] "Lincoln Continental" "Chrysler Imperial"
4 [5] "Camaro Z28" "Ford Pantera L"
5 [7] "Maserati Bora"
6 > rownames(auta)[auta$gear == 5]
7 [1] "Porsche 914-2" "Lotus Europa" "Ford Pantera L"
8 [4] "Ferrari Dino" "Maserati Bora"
9 >
10 > rownames(auta)[auta$hp > 200 & auta$gear == 5]
11 [1] "Ford Pantera L" "Maserati Bora"
12 > auta[auta$hp > 200 & auta$gear == 5, ]
13      mpg   cyl  disp    hp  drat    wt  qsec vs am gear carb
14 Ford Pantera L 15.8   8  351 264 4.22 3.17 14.5  0  1   5    4
15 Maserati Bora  15.0   8  301 335 3.54 3.57 14.6  0  1   5    8
16      disp_group
17 Ford Pantera L   nad 300
18 Maserati Bora    nad 300
```

Príklad 11

Nájdite tie automobily, ktorých zdvihový objem motora je menej ako 100 kubických palcov, alebo majú dva karburátory.

Riešenie 11

```
1 > rownames(auta)[auta$disp < 100]
2 [1] "Fiat 128" "Honda Civic" "Toyota Corolla"
3 [4] "Fiat X1-9" "Lotus Europa"
4 > rownames(auta)[auta$disp_group == "do 100"]
5 [1] "Fiat 128" "Honda Civic" "Toyota Corolla"
6 [4] "Fiat X1-9" "Lotus Europa"
7 > rownames(auta)[auta$carb == 2]
8 [1] "Hornet Sportabout" "Merc 240D" "Merc 230"
9 [4] "Honda Civic" "Dodge Challenger" "AMC Javelin"
10 [7] "Pontiac Firebird" "Porsche 914-2" "Lotus Europa"
```

```

11 [10] "Volvo 142E"
12 >
13 > rownames(auta)[auta$disp < 100 | auta$carb == 2]
14 [1] "Hornet Sportabout" "Merc 240D" "Merc 230"
15 [4] "Fiat 128" "Honda Civic" "Toyota Corolla"
16 [7] "Dodge Challenger" "AMC Javelin" "Pontiac Firebird"
17 [10] "Fiat X1-9" "Porsche 914-2" "Lotus Europa"
18 [13] "Volvo 142E"
19 > auta[auta$disp < 100 | auta$carb == 2, ]
20      mpg   cyl  disp    hp  drat    wt   qsec    vs  am gear
21 Hornet Sportabout 18.7   8 360.0  175  3.15  3.440  17.02   0   0    3
22 Merc 240D        24.4   4 146.7   62  3.69  3.190  20.00   1   0    4
23 Merc 230         22.8   4 140.8   95  3.92  3.150  22.90   1   0    4
24 Fiat 128         32.4   4  78.7   66  4.08  2.200  19.47   1   1    4
25 Honda Civic      30.4   4  75.7   52  4.93  1.615  18.52   1   1    4
26 Toyota Corolla   33.9   4  71.1   65  4.22  1.835  19.90   1   1    4
27 Dodge Challenger 15.5   8 318.0  150  2.76  3.520  16.87   0   0    3
28 AMC Javelin      15.2   8 304.0  150  3.15  3.435  17.30   0   0    3
29 Pontiac Firebird 19.2   8 400.0  175  3.08  3.845  17.05   0   0    3
30 Fiat X1-9        27.3   4  79.0   66  4.08  1.935  18.90   1   1    4
31 Porsche 914-2    26.0   4 120.3   91  4.43  2.140  16.70   0   1    5
32 Lotus Europa     30.4   4  95.1  113  3.77  1.513  16.90   1   1    5
33 Volvo 142E       21.4   4 121.0  109  4.11  2.780  18.60   1   1    4
34      carb disp_group
35 Hornet Sportabout    2    nad 300
36 Merc 240D           2    100-200
37 Merc 230            2    100-200
38 Fiat 128            1      do 100
39 Honda Civic         2      do 100
40 Toyota Corolla      1      do 100
41 Dodge Challenger    2    nad 300
42 AMC Javelin         2    nad 300
43 Pontiac Firebird    2    nad 300
44 Fiat X1-9           1      do 100
45 Porsche 914-2       2    100-200
46 Lotus Europa        2      do 100
47 Volvo 142E          2    100-200

```

Pre efektívnejšiu prácu s údajmi odporúčame používať funkcie z balíka `dplyr`, ktorý ponúka oveľa prepracovanejšie možnosti manipulácie dát (výber, filtrovanie, triedenie, spájanie či vytváranie nových premenných a podobne). Viac informácií o práci s využitím tohto balíka nájdete v pokračovaní tejto knihy (*R* *snadno a rýchle 2 Vizualizace dat a programování*).

Zoznam použitej literatúry

- [1] ADLER J., 2010. *R in a Nutshell – a Desktop Quick Reference*. O'Reilly, ISBN 978-0-596-80170-0.
- [2] DATAFLAIR TEAM. *Data Analytics Tutorial for Beginners – From Beginner to Pro in 10 Mins!* [online]. 2019 [cit. 2020-03-04].
Dostupné na: <https://data-flair.training/blogs/data-analytics-tutorial/>
- [3] HAFNER, S. *An Introduction to R for Beginners* [online]. 2019 [cit.2020-01-20]. Dostupné na: https://www.researchgate.net/publication/325170649_An_Introduction_to_R_for_Beginners
- [4] KOECH, H. *11 Best R Programming IDE and editors* [online]. 2019 [cit. 2020-03-15].
Dostupné na: <https://www.dunebook.com/best-r-programming-ide/>
- [5] KVASNÍČKA, M. *Analýza a vizualizace ekonomických dat* [online]. 2017 [cit. 2020-01-17]. Dostupné na: https://is.muni.cz/el/1456/podzim2016/BPE_AVED/um/70030968/
- [6] R CORE TEAM. *A language and environment for statistical computing* [online]. R Foundation for Statistical Computing, Vienna, Austria, 2017 [cit. 2020-01-17]. Dostupné na: <https://www.R-project.org/>
- [7] RSTUDIO TEAM. *RStudio: Integrated Development for R* [online]. RStudio, Inc., Boston, MA, 2020 [cit. 2020-01-25]. Dostupné na: <http://www.rstudio.com/>
- [8] LONG, J. D., TEETOR, P. *R Cookbook: Proven Recipes for Data Analysis, Statistics, and Graphics*. O'Reilly Media; 2nd edition 2019. ISBN 1492040681. Dostupné tiež na: <https://rc2e.com/>
- [9] TORFS, P. J. J. F. BREUER, C. *A (very) short introduction to R* [online]. 2012 [cit. 2020-02-01]. Dostupné na: <https://cran.r-project.org/doc/contrib/Torfs+Brauer-Short-R-Intro.pdf>
- [10] VENABLES, W. N. SMITH, D. M. *An Introduction to R. Notes on R: A Programming Environment for Data Analysis and Graphics* [online]. 2020 [cit. 2020-04-14].
Dostupné na: <https://cran.r-project.org/doc/manuals/r-release/R-intro.pdf>

-
- [11] ZUUR, Alain F., Elena N. IENO a H. W. G. MEESTERS, 2009. *A beginner's guide to R*. Dordrecht: Springer. Use R!. ISBN 978-0-387-93836-3.

Názov	R snadno a rychle 1 Úvod do jazyka
Autori	Ing. Jakub Danko, PhD. Ing. Karel Šafr, Ph.D.
Vydavateľ	Vysoká škola ekonomická v Praze Nakladatelství Oeconomica
Vydanie	Prvé vydanie v elektronickej forme
Redakčná úprava	Mgr. Ludmila Doudová
Grafický návrh	Daniel Hamerník, DiS.
Počet strán	161
DTP	Vysoká škola ekonomická v Praze Nakladatelství Oeconomica
Odporúčaná cena	zadarmo

ISBN 978-80-245-2380-4