

Vysoká škola ekonomická v Praze

R SNADNO A RYCHLE 2

VIZUALIZACE DAT A PROGRAMOVÁNÍ



Karel Šafr, Jakub Danko

2020



OECONOMICA

Nakladatelství VŠE

Autoři:

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

Ing. Jakub Danko, 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)

Poděkování

Na tomto místě bychom rádi poděkovali zejména doc. RNDr. Ivaně Malé, CSc. (Vysoká škola ekonomická v Praze) za velké množství času, který nám věnovala při psaní této naší prvotiny. Bez jejích odborných připomínek, komentářů a dotazů by zcela jistě tato kniha nevznikla ve stávajícím formátu, rozsahu a kvalitě. Dále bychom chtěli poděkovat recenzentům této knihy, kterými jsou: Mgr. Vítězslav Štembera, Ph.D. (Technická univerzita Vídeň) a Ing. Zdeněk Šulc, Ph.D. (Vysoká škola ekonomická v Praze), za množství a podnětnost připomínek, které oba recenzenti vznesli, a velmi si vážíme toho, že věnovali tolik svého volného času detailnímu recenzování této knihy. Dále bychom rádi poděkovali Ing. Veronice Ptáčkové, MUDr. Leoně Karáskové, Květě a Nele Šafrové, které provedly prvotní gramatickou a slohovou korekturu této knihy. V neposlední řadě patří velký dík našim nejbližším a rodinám, které nás podporovaly a motivovaly při psaní této knihy.

V neposlední řadě je nutno zmínit, že vznik knihy byl podpořen projektem interní rozvojové soutěže VŠE v Praze (IRS) F4/18/2019.

Autoři

Obsah

Úvod	6
1. Cykly	7
1.1 For cyklus	9
1.2 While cyklus	17
2. Podmínky	22
2.1 Podmínka if	22
2.2 Větvení podmínek pomocí else	24
2.3 Větvení podmínek pomocí else if	25
2.4 Větvení podmínek ifelse	28
3. Příkazy next, break a repeat cyklus	30
3.1 Příkaz next a break	30
3.2 Cyklus repeat	32
4. Funkce apply, by a aggregate	34
4.1 Funkce apply	34
4.2 Funkce lapply, sapply, vapply, replicate	43
4.2.1 Funkce lapply	43
4.2.2 Funkce sapply	46
4.2.3 Funkce vapply	48
4.2.4 Funkce replicate	49
4.3 Funkce mapply	50
4.4 Funkce rapply	52
4.5 Funkce tapply a by funkce	54
4.6 Funkce aggregate a ts objekt	58
5. Tvorba vlastních funkcí	64
5.1 Uživatelem definované funkce	64
5.1.1 Jméno funkce	65
5.1.2 Parametry funkce	66
5.1.3 Tělo funkce	68
5.1.4 Návrátová hodnota funkce	72
5.2 Anonymní funkce	74
5.3 Rekurzivní funkce	75

6. Základy tvorby grafů v jazyce R	78
6.1 Úvodem ke grafům v R	78
6.2 Base graphics	79
6.2.1 Funkce plot()	80
6.2.2 Základní editace grafu a barvy	83
6.2.3 Základní typy grafů	90
6.2.4 Funkce par() a ostatní nízkoúrovňové funkce	107
6.2.5 Tvorba legendy	116
6.2.6 Práce s grafickými okny v softwaru R	119
6.2.7 Jak grafy ukládat	120
6.3 Ggplot2	122
6.3.1 Základní model grafu	122
6.3.2 Grafická reprezentace dat – geom a stats	127
6.3.3 Funkce scale	135
6.3.4 Nadpisy a popisky os	139
6.3.5 Více grafů v jedné grafické oblasti	143
6.3.6 Ukládání grafu	147
7. Funkce merge a knihovna dplyr	149
7.1 Funkce merge	149
7.2 Knihovna dplyr	153
7.2.1 Pipeline operátor	153
7.2.2 Filtrování pozorování	154
7.2.3 Seřídění datasetu	155
7.2.4 Filtrování proměnných	156
7.2.5 Tvorba nové proměnné	158
7.2.6 Statistické shrnutí dat	159
7.2.7 Práce se skupinami dat	160
7.2.8 Spojování datasetů pomocí balíčku dplyr	160
8. Formátování kódu	163
8.1 Jména	164
8.1.1 Jména objektů a funkcí	164
8.1.2 Jména souborů	164
8.2 Obecná syntaxe	165
8.2.1 Psaní mezer	165
8.2.2 Logické operátory	166
8.2.3 Uvozovky	167
8.3 Bloky kódu a delší zápis kódu	167
Seznam obrázků	169
Seznam tabulek	171
Literatura	172

Úvod

Vážené čtenářky, vážení čtenáři,

dostal se vám do rukou druhý díl knihy „R snadno a rychle“. V prvním díle knihy „R snadno a rychle 1 Úvod do jazyka“ jsme představili základy jazyka R. Primárně jsme se zaměřili na to, jak definovat objekty, jaké mají vlastnosti, a vysvětlili jsme si základy práce v programovacím jazyku R. V knize, kterou jste nyní otevřeli, se již nebudeme zabývat základy, jako je i načítání dat, ale rovnou se budeme věnovat základům programování a vizualizace dat v R. Musíme vás tedy upozornit na to, že kniha se již nezabývá základy práce v jazyku R – definice objektů, import dat, operátor nebo základní funkce. Předpokládáme, že uživatelé zmíněné znalosti bez problémů ovládají. Kniha je tak vhodná pro pokročilé začátečníky, ať už mezi studenty, nebo veřejností, která si chce rozšířit znalosti a usnadnit si tak výpočty v pracovním životě.

V první kapitole vám představíme cykly, které jsou důležité zejména z pohledu automatizace zpracování řešené úlohy. V návaznosti na cykly si ukážeme, jak fungují podmínky. Pomocí nich totiž dokážeme kód tzv. větvit – podmínit chování v určité situaci. Uvedeme příklad: chceme, aby se pro sudá čísla realizoval jiný výpočet než pro lichá. Kapitola „Next, break a repeat cyklus“ uzavře dvě zmíněné kapitoly představením tzv. obecného cyklu „repeat“. V navazující kapitole („Funkce apply, by a aggregate“) vám vysvětlíme, jak fungují velmi oblíbené apply funkce a s nimi spojená agregace dat. Další kapitolou je „Tvorba vlastních funkcí“, která představí, jak lze v programovacím jazyku R vytvořit vlastní uživatelskou funkci. Kapitola „Základy tvorby grafů v jazyce R“ se věnuje tvorbě grafů pomocí základní knihovny a populární knihovny ggplot2. V předposlední kapitole „Funkce merge a knihovna dplyr“ si ukážeme, jak data spojit, agregovat nebo filtrovat. Poslední kapitola se zaměří na formátování kódu jazyka R.

Přejeme vám mnoho úspěchů při programování a analýze dat pomocí jazyka R. Doufáme, že se vám bude kniha líbit a usnadní vám práci v jazyce R.

Hodně štěstí!

Tým autorů

Kapitola 1

Cykly

V rámci předchozích kapitol jsme si představili základy práce s jazykem R. Ukázali jsme si, jak se pracuje s jednotlivými objekty, představili jsme si základní funkce a vysvětlili základy práce zpracování dat. V rámci navazujících kapitol si představíme pokročilejší programovací techniky. Výhodou programovacího jazyka R oproti statistickým softwarům je zejména automatizace, tvorba vlastních funkcí, grafické knihovny a pokročilejší správa dat. V této kapitole si představíme metody pro automatické opakování vykonávaného kódu neboli cykly.

Cyklus (neboli „smyčka“) je cyklicky se opakující sekvence kódu. Cyklicky se opakující část kódu se nachází uvnitř složených závorek za definicí funkcí cyklu. Cyklus je vykonán pro předem definovanou sekvenci hodnot, má stanovený počet opakování (**for** cyklus), anebo pokud je splněná podmínka, která povoluje opakování cyklu (**while** cyklus). Jedna realizace cyklu se nazývá iterací. Iteraci lze též v češtině nazvat průchod cyklem.

Cykly jsou součástí většiny programovacích jazyků a jedná se o jejich nejdůležitější komponentu. Výhodou cyklů je zejména automatizace procesů a zjednodušení zápisu kódu. Například postupné zobrazení hodnot od 1 do 10 bychom bez cyklu mohli zapsat jako postupné zadáváním příkazu `print()`.

```
1 > print(1)
2 [1] 1
3 > print(2)
4 [1] 2
5 > print(3)
6 [1] 3
7 > print(4)
8 [1] 4
9 > print(5)
10 [1] 5
11 > print(6)
```

```
12 [1] 6
13 > print(7)
14 [1] 7
15 > print(8)
16 [1] 8
17 > print(9)
18 [1] 9
19 > print(10)
20 [1] 10
```

Pomocí cyklu si můžeme výrazně usnadnit zápis kódu na jeden řádek. Stačí, když použijeme for cyklus a funkci `print()` vložíme do těla cyklu.

```
1 > for (i in 1:10) {print(i)}
2 [1] 1
3 [1] 2
4 [1] 3
5 [1] 4
6 [1] 5
7 [1] 6
8 [1] 7
9 [1] 8
10 [1] 9
11 [1] 10
```

Ne v každé situaci je nutné použít cyklus. Pokud by nám šlo o čisté zobrazení hodnot od jedné do desíti, můžeme použít funkci `seq()`. Funkce `seq()` vytváří sekvenci hodnot od definované první hodnoty (parametr `from`) do poslední hodnoty (parametr `to`). Délku sekvence můžeme ovlivnit pomocí parametru `length.out` nebo `by`. Parametr `length.out` specifikuje délku výsledného vektoru (velikost kroku si vypočítá funkce sama), anebo můžeme použít parametr `by`, který specifikuje délku kroku mezi jednotlivými prvky vektoru (a délka vektoru je výsledně získána samotnou funkcí).

```
1 > seq(from = 2, to = 8, length.out = 4)
2 [1] 2 4 6 8
3 > seq(from = 2, to = 8, by = 1.5)
4 [1] 2.0 3.5 5.0 6.5 8.0
```

Cyklus najde své použití zejména v situaci, kdy chceme v rámci iterace provést více úloh, jež nelze provést pro celý vektor/objekt najednou.

1.1 For cyklus

Cyklus `for` je cyklus, u kterého je předem známý počet opakování. Cyklus je vykonáván pomocí tzv. iterátoru (iterační proměnná, popř. řídicí proměnná), který postupně nabývá hodnot z předem definované množiny. Množinou hodnot se rozumí obvykle prvky atomického vektoru hodnot, ale je možné iterovat i přes prvky listu. Prozatím se ale omezíme na atomický vektor hodnot.

Cyklus vytvoříme pomocí příkazu `for`. V kulatých závorkách za tímto příkazem je nutné definovat iterační proměnnou. Iterační proměnnou definujeme tak, že specifikujeme její jméno a dále její rozsah. Tedy například výraz `i in 1:10` způsobí, že iterační proměnná `i` bude nabývat hodnot od jedné do desíti v jednotlivých průchodech cyklem. Samotné tělo cyklu definujeme až následně. Standardně je tělo cyklu uzavřeno do bloku složených závorek (`{...}`). Tělo cyklu obsahuje kód, který se má v jednotlivých iteracích vykonat.

```
1 for (# definice iteracni promene) {  
2 # telo cyklu  
3 }
```

V předchozí kapitole jsme si ukázali `for` cyklus, kde se v každé iteraci zobrazila hodnota iterátoru. Cyklus lze ale rozepsat i na více řádek tak, jak ilustruje následující kód.

```
1 > for (i in 1:10) {  
2 +   print(i)  
3 + }  
4 [1] 1  
5 [1] 2  
6 [1] 3  
7 [1] 4  
8 [1] 5  
9 [1] 6  
10 [1] 7  
11 [1] 8  
12 [1] 9  
13 [1] 10
```

Cyklus `for` lze zapsat i ve zkrácené formě. Jedná se o situaci, kdy lze kód napsat na jeden řádek. V tu chvíli můžeme vynechat složené závorky. Je nutné si ale uvědomit, že se iteruje pouze kód do konce řádku či do prvního středníku na daném řádku.

```
1 > for (i in 1:10) print(i)  
2 [1] 1  
3 [1] 2  
4 [1] 3  
5 [1] 4  
6 [1] 5
```

```
7 [1] 6
8 [1] 7
9 [1] 8
10 [1] 9
11 [1] 10
```

Alternativně lze zapsat cykly i pomocí odsazení (bez složených závorek). Iteruje se zde pouze první odsazená řádka za definicí `for` cyklu.

```
1 > for (i in 1:10)
2 +   print(i)
3 [1] 1
4 [1] 2
5 [1] 3
6 [1] 4
7 [1] 5
8 [1] 6
9 [1] 7
10 [1] 8
11 [1] 9
12 [1] 10
13 >   print("ahoj")
14 [1] "ahoj "
```

Klíčovou roli v cyklu zastává řídicí proměnná (iterační). Jak již bylo řečeno, řídicí proměnná nabývá postupně hodnot z množiny, přes jejíž prvky iteruje. Tato řídicí proměnná nám může usnadnit přístup k dílčím prvkům jiných objektů. Toto si ukážeme na následujícím příkladu, kde vypíšeme jednotlivé sloupce matice `X`. Nejprve si ale matici `i` zobrazíme, abychom poznali rozdíl od výsledného výstupu z cyklu (pomocí funkce `print()`).

```
1 > # Definice matice:
2 > X <- matrix(1:9, 3)
3 > print(X)
4      [,1] [,2] [,3]
5 [1,]    1    4    7
6 [2,]    2    5    8
7 [3,]    3    6    9
8 >
9 > # FOR cyklus:
10 > for (i in 1:ncol(X)) {
11 +   print(X[, i])
12 + }
13 [1] 1 2 3
14 [1] 4 5 6
15 [1] 7 8 9
```

Tento kód provede iterace od jedné do počtu sloupců proměnné (X) a zobrazí hodnoty v těchto sloupcích. Co bychom ale měli udělat, kdybychom chtěli vytvořit cyklus, který bude postupně zobrazovat jednotlivé prvky matice? Máme dvě základní možnosti: 1) převést matici X na vektor a postupně její hodnoty zobrazit, 2) vnořit dva cykly do sebe a postupovat po řádcích nebo po sloupcích. Převedení matice na vektor a iterace od délky daného vektoru (length) je velmi snadná úloha, pojďme si tedy ukázat vnoření dvou cyklů do sebe.

```
1 > # Definice matice
2 > X <- matrix(1:9, 3)
3 > print(X)
4      [,1] [,2] [,3]
5 [1,]    1    4    7
6 [2,]    2    5    8
7 [3,]    3    6    9
8 >
9 > # Cyklus
10 > for (i in 1:ncol(X)) {
11 +   #Zobrazení informace o daném cyklu:
12 +   print(paste("Vnější cyklus:",i))
13 +   for (j in 1:nrow(X)) {
14 +     print(X[i, j])
15 +   }
16 + }
17 [1] "Vnější cyklus: 1"
18 [1] 1
19 [1] 4
20 [1] 7
21 [1] "Vnější cyklus: 2"
22 [1] 2
23 [1] 5
24 [1] 8
25 [1] "Vnější cyklus: 3"
26 [1] 3
27 [1] 6
28 [1] 9
```

Před samotný výkon vnitřního cyklu (s iterátorem j) je vložena funkce `print()`, která nás informuje o tom, jaké hodnoty nabývá iterátor (i) vnějšího cyklu. Cyklus tedy postupuje tak, že nejprve vnější cyklus přiřadí do iteráčnické proměnné i hodnotu 1 a následně vnitřní cyklus postupně dosazuje do iteráčnické proměnné j hodnoty 1 až 3. Tedy první hodnoty, které cyklus zobrazí, jsou X[1, 1], X[1, 2], X[1, 3]. Tyto hodnoty odpovídají první iteraci vnějšího cyklu a jedné až třetí iteraci vnitřního cyklu. Následně cyklus postoupí do druhého kola vnějšího for cyklu a zobrazí: X[2, 1], X[2, 2], X[2, 3]. Tedy v druhém kole (kdy iteráčnická proměnná vnějšího cyklu i je rovna 2) znovu vnitřní cyk-

lus nabude hodnot 1, 2 a 3 (iterační proměnná *j*). V poslední fázi vnějšího cyklu (kdy *i* je rovna 3) budou zobrazeny prvky: *X*[3, 1], *X*[3, 2], *X*[3, 3]. Samozřejmě v případě, kdy si vytváříme vlastní cyklus, není vůbec třeba vypisovat informaci o tom, jaké hodnoty iterační proměnná v určité fázi cyklu nabývá. V našem případě nabývá hodnot `print(paste("Vnější cyklus:", i))`. Zde tuto informaci uvádíme pro zpřehlednění výkonu jednotlivých iterací.

Všimněte si, že pro každý cyklus využíváme jinak pojmenovanou řídicí proměnnou. To je z toho důvodu, že řídicí proměnná se po dobu realizace těla cyklu chová jako jakýkoliv jiný objekt v **R**. Teprve až po dokončení kola iterace nabude další hodnoty z přiřazené množiny hodnot. Pokud jste před realizací cyklu vytvořili v **R** proměnnou stejného jména, tak ji cyklus přepíše. Po dokončení cyklu bude nabývat poslední hodnotu přiřazené množiny, přes jejíž prvky iteruje.

Zajímavá situace nastává v případě, kdy vnoříme do sebe dva cykly a oba budou používat stejně pojmenovanou řídicí proměnnou. Vnější cyklus přiřadí do řídicí proměnné první prvek vnější množiny, avšak vnitřní cyklus bude postupně do dané proměnné přiřazovat prvky své (druhé) množiny. Na konci první iterace vnějšího cyklu bude v řídicí proměnné poslední hodnota množiny vnitřního cyklu. Na začátku druhé iterace vnějšího cyklu ale bude již v iterátoru zase prvek (druhé) množiny vnějšího cyklu. A takto se situace bude opakovat až do konce obou cyklů.

```
1 > for (i in 1:3){
2 +   print(paste("zacatek vnejsiho cyklu", i))
3 +
4 +   for (i in 4:5) {
5 +     print(paste("vnitrni cyklus", i))
6 +   }
7 +
8 +   print(paste("konec vnejsiho cyklu", i))
9 + }
10
11 [1] "zacatek vnejsiho cyklu 1"
12 [1] "vnitrni cyklus 4"
13 [1] "vnitrni cyklus 5"
14 [1] "konec vnejsiho cyklu 1"
15 [1] "zacatek vnejsiho cyklu 2"
16 [1] "vnitrni cyklus 4"
17 [1] "vnitrni cyklus 5"
18 [1] "konec vnejsiho cyklu 2"
19 [1] "zacatek vnejsiho cyklu 3"
20 [1] "vnitrni cyklus 4"
21 [1] "vnitrni cyklus 5"
22 [1] "konec vnejsiho cyklu 3"
```

Podobná situace bude nastávat i v případě, kdy se rozhodneme předefinovat samotnou iterační proměnnou v samotném běhu cyklu.

```

1 > for (i in 1:3){
2 +   print(i)
3 +   i=1
4 +   print(i)
5 + }
6 [1] 1
7 [1] 1
8 [1] 2
9 [1] 1
10 [1] 3
11 [1] 1

```

Zde je ale nutné upozornit čtenáře, že kód, kde do iterační proměnné jsou přiřazované jakékoliv hodnoty v rámci těla cyklu, je vyloženě nesprávný a uživatel by se měl vyvarovat tomuto použití. Kód se takovouto aplikací znepráhlední a často nastávají neočekávané (a chybné) situace.

Řekli jsme si, že jednotlivé elementy for cyklu jsou vykonávané pro jednotlivé prvky množiny hodnot (obvykle vektoru), přes který má řídicí proměnná iterovat. Cyklus nemusí být proveden pouze přes atomický vektor hodnot, ale lze jej realizovat i přes list (list je přece také vektor), anebo přes jakékoliv jiné objekty, k jejichž vnitřním datům lze přistupovat přes předem definovanou dimenzi (např. `data.frame`). Obvykle se jedná o objekty, u nichž lze zjistit délku, například pomocí funkce `length`, a přistoupit na této délce k částečné výseči daného objektu.

```

1 > # Definice listu:
2 > l <- list(c = 1,
3 +         b = c(1, 3),
4 +         c = matrix(1:9, 3))
5 > # Zobrazíme si objekt aby jsme vedeli přes co iterujeme
6 > print(l)
7 $c
8 [1] 1
9
10 $b
11 [1] 1 3
12
13 $c
14      [,1] [,2] [,3]
15 [1,]    1    4    7
16 [2,]    2    5    8
17 [3,]    3    6    9
18
19 > # Samotné provedení for cyklu:
20 > for (i in l) {

```

```

21 +   print(i)
22 + }
23 [1] 1
24 [1] 1 3
25      [,1] [,2] [,3]
26 [1,]    1    4    7
27 [2,]    2    5    8
28 [3,]    3    6    9

```

Jak si zcela jistě pamatujete z úvodních kapitol knihy, tak objekt typu `list` je vlastně formou vektoru. Nejprve se tedy zobrazil podlist `$a`, pak `$b` a následně `$c`.

Iterovat můžeme i přes objekt typu `data.frame`. Takové iterování si ukážeme na objektu `mtcars`. Do iterační proměnné jsou v jednotlivých iteracích přiřazované sloupce (proměnné) z datasetu `mtcars`. Délka objektu je zde čerpána z počtu proměnných obsažených v objektu typu `data.frame`, tedy odpovídá počtu sloupců.

```

1 > for (i in mtcars) {
2 +   print(mean(i))
3 + }
4 [1] 20.09062
5 [1] 6.1875
6 [1] 230.7219
7 [1] 146.6875
8 [1] 3.596563
9 [1] 3.21725
10 [1] 17.84875
11 [1] 0.4375
12 [1] 0.40625
13 [1] 3.6875
14 [1] 2.8125

```

Analogický výstup bychom získali i použitím příkazu `colMeans(mtcars)`.

Většina ostatních objektů ale nemá parametr délky (atribut `length`). Například matice má atributy dimenzí, ale ne délky. Kdybychom se ale i tak rozhodli iterovat přes matici (přes sloupce matice), tak získáme nesmyslný výstup.

```

1 > X <- matrix(1:4, 2)
2 > print(X)
3      [,1] [,2]
4 [1,]    1    3
5 [2,]    2    4
6 > for (i in X) {
7 +   print(X)
8 + }
9      [,1] [,2]

```

```

10 [1,]      1      3
11 [2,]      2      4
12      [,1] [,2]
13 [1,]      1      3
14 [2,]      2      4
15      [,1] [,2]
16 [1,]      1      3
17 [2,]      2      4
18      [,1] [,2]
19 [1,]      1      3
20 [2,]      2      4

```

Program **R** není schopen přiřadit iterátoru jednotlivé elementy objektu (řádky, sloupce či jednu buňka matice), a tak se provede iterace přes součet dimenzí, ale v každé iteraci nabývá iterátor hodnoty samotné celé matice. Pokud se tedy uživatel rozhodne iterovat přes celý objekt, měl by se předem ujistit o třídě a dimenzích takového objektu.

Obykle se v programovacím jazyku **R** volí atomický vektor hodnot, jehož hodnoty iterátor nabývá. Hlavním důvodem je přehlednost, pak i to, že **R** má rodinu `apply()` funkcí, která usnadňuje vykonání cyklické operace přes určité druhy objektů.

Příklad 1

Spočítejte pomocí `for` cyklu variační koeficient pro proměnné datasetu `mtcars`.

Řešení 1

```

1 > for (i in mtcars) {
2 +   print(sd(i) / mean(i))
3 + }
4 [1] 0.2999881
5 [1] 0.2886338
6 [1] 0.5371779
7 [1] 0.4674077
8 [1] 0.1486638
9 [1] 0.3041285
10 [1] 0.1001159
11 [1] 1.152037
12 [1] 1.228285
13 [1] 0.2000825
14 [1] 0.5742933

```

Příklad 2

Máme vektor $x=1:10$ a chceme pomocí cyklu vytvořit nový vektor (k), který bude obsahovat kumulace vektoru x . Tedy kumulovanou sumu hodnot vektoru x od prvního prvku do pozice daného prvku, pro kterou sumu provádíme.

Řešení 2

```
1 > x <- 1:10
2 > print(x)
3 [1] 1 2 3 4 5 6 7 8 9 10
4 > # Prazdny vektor k:
5 > k <- c()
6 > # Cyklus
7 > for (i in 1:length(x)) {
8 +   k[i] <- sum(x[1:i])
9 + }
10 > # Vysledek:
11 > print(k)
12 [1] 1 3 6 10 15 21 28 36 45 55
```

Příklad 3

Vytvořte matici X , která bude mít pět řádek a čtyři sloupce. Následně pomocí `for` cyklu(ů) umístěte do prvků této matice indexy daného prvku matice (buňky).

Řešení 3

```
1 > # Vytvorime si prve matici:
2 > X <- matrix(ncol=4, nrow=5)
3 > print(X)
4      [,1] [,2] [,3] [,4]
5 [1,]  NA  NA  NA  NA
6 [2,]  NA  NA  NA  NA
7 [3,]  NA  NA  NA  NA
8 [4,]  NA  NA  NA  NA
9 [5,]  NA  NA  NA  NA
10 > # Nasledne pomoci dvou cyklu projdeme vsechny prvky matice:
11 > for (i in 1:nrow(X)){
12 +   for (j in 1:ncol(X)){
13 +     X[i, j]<-paste0(i, ",", j)
14 +   }
15 + }
```



```

16 > print(X)
17      [,1]  [,2]  [,3]  [,4]
18 [1,] "1, 1" "1, 2" "1, 3" "1, 4"
19 [2,] "2, 1" "2, 2" "2, 3" "2, 4"
20 [3,] "3, 1" "3, 2" "3, 3" "3, 4"
21 [4,] "4, 1" "4, 2" "4, 3" "4, 4"
22 [5,] "5, 1" "5, 2" "5, 3" "5, 4"

```

1.2 While cyklus

While cyklus vykonává (opakuje) tělo cyklu do té doby, do které je splněná podmínka pro výkon cyklu. Výsledek této podmínky musí být interpretovatelný jako logická proměnná. Tedy i čísla 0 a 1 lze použít jako vstup do podmínky cyklu. Standardně se ale jedná o logický výraz, který testuje nějakou vlastnost. Následující kód ilustruje použití while cyklu.

```

1 while (# podminka) {
2   # telo cyklu
3 }

```

V případě, kdy je podmínka vyhodnocena jako FALSE, tak další kolo cyklu již není vykonáno, a to i v případě, že doposud nebyla žádná iterace vykonána vůbec.

Na následujícím příkladu je ilustrována situace použití while cyklu, která je analogická k for cyklu (`for (i in 1:10)`). Tedy v jednotlivých kolech iterace vypíšeme hodnotu proměnné `i` od jedné do desíti.

```

1 > i = 1
2 > while (i <= 10) {
3   +   print(i)
4   +   i = i + 1
5   + }
6 [1] 1
7 [1] 2
8 [1] 3
9 [1] 4
10 [1] 5
11 [1] 6
12 [1] 7
13 [1] 8
14 [1] 9
15 [1] 10

```

Výraz v kulaté závorce je logickým testem toho, zda je `i` menší nebo rovno deseti. Pokud

výsledkem tohoto testu je hodnota `TRUE`, vykoná se kód ve složené závorce. Pokud by však výsledek jednou nabyl hodnoty `FALSE`, tak se cyklus úplně zastaví a žádná další iterace se již nevykoná.

V našem případě proměnná `i` plní funkci iterátoru, a tak je celý cyklus velmi podobný `for` cyklu. Na tomto příkladu je hezky vidět, proč je `for` cyklus považován za formu `while` cyklu. V programovacím jazyku **R** je ale situace o něco složitější, protože je zde ještě obecnější forma cyklu `while`. Jedná se o cyklus jménem `repeat`. Na tento cyklus se ale podíváme až v následujících kapitolách.

Důležité je zde upozornit, že proměnná `i` musí být zavedená ještě před samotným zápisem cyklu, jinak by cyklus navrátil `error (missing value)`. Protože i pro první spuštění cyklu dochází k testování podmínky toho, zda se má cyklus spustit! Při definici podmínky musí tedy uživatel brát tuto situaci předem v potaz. Stává se, že mnoho lidí definuje cyklus, jehož výchozí podmínka nemůže být pro první kolo iterace nikdy splněná. Podmínka samozřejmě nemusí být jen jedna a logické výrazy můžeme kombinovat.

```
1 > i = 1
2 > j = 1
3 > while (i >= 1 & j <= 1 | abs(i * j) < 10) {
4 +   i = i * j + 1
5 +   j = j * i - 1
6 +   print(paste("i:", i, ", j:", j, "abs(i*j):", abs(i * j)))
7 + }
8 [1] "i: 2 , j: 1 abs(i*j): 2"
9 [1] "i: 3 , j: 2 abs(i*j): 6"
10 [1] "i: 7 , j: 13 abs(i*j): 91"
```

V každé iteraci cyklu je proveden předem definovaný výpočet a následně se zobrazí hodnoty jednotlivých proměnných, které se budou před začátkem dalšího cyklu ověřovat v podmínce. Pro zobrazení proměnných využíváme stejné funkce jako minule v předcházejícím příkladu – funkce `print()`.

V praxi se může stát, že nebudeme vědět, zda vůbec nastane situace, kdy a jestli bude podmínka nesplněna (`FALSE`). Tedy nejsme schopni s jistotou říct, zda se cyklus zastaví. Tato situace je velmi podobná tomu, kdy by cyklus splnil podmínku až po velkém množství iterací, a chceme se předem ujistit, že tato situace nenastane. Tedy chceme zastavit cyklus po předem daném počtu iterací.

Tyto dva problémy lze vyřešit dvěma způsoby. Buď musíme očistit tuto situaci pomocí řídicích procedur `NEXT` a `BREAK` v kombinaci s podmínkami typu `if`, nebo přidat do cyklu `while` iterátor, který předem definuje maximální počet iterací. `NEXT` a `BREAK` příkazy probereme až v následujících kapitolách. Teď si ukážeme řešení této situace pro příklad, kdy vytvoříme nekonečný cyklus ošetřený pomocí iterátoru.

```
1 > i = 0
2 > x = 1
```

```

3 > while (x != 2 & i < 100) {
4 +   x = x + 0.13
5 +   i = i + 1
6 + }
7 > print(i)
8 [1] 100

```

Příklad je nadefinovaný tak, že nikdy nebude splněná podmínka o rovnosti x dvou. Z tohoto důvodu je nutné přidat proměnnou iterátoru jako součet podmínky – pro ošetření celkového počtu cyklů. Nezapomeňte inkrementovat iterátor, oproti for cyklu to za vás while cyklus neudělá.

Přesto, že se vám může zdát předchozí příklad relativně nesmyslný (kdo by definoval ukončení cyklu na určitou nerovnost, která nenastává), tak je nutné vědět, že nesmíme zapomenout na to, že **R** má přesnost výpočtu na šestnáct desetinných míst ($1e-16$!). Proto mnoho lidí o hodnotách, které jsou menší než $1e-16$, mluví jako o technicky nulových. Může se pak u některých typů výpočtů snadno stát, že analytickým řešením je nulový výsledek, ale technicky získáme číslo například $1e-18$. Tedy číslo, které je nenulové, ale je již v pásmu zaokrouhlovací chyby softwaru. Pokud by ale while cyklus testoval toto číslo na rovnost nule, nedošlo by k ukončení cyklu, protože číslo je pouze technicky nulové.

Příklad 4

Pomocí while cyklu zobrazte hodnoty proměnné x (definované sekvencí hodnot 10:1) a v případě, kdy bude hodnota menší než 5, tak cyklus zastavte.

Řešení 4

```

1 > #Definice iteratoru
2 > i = 1
3 > #while cyklus se dvema podmínkami
4 > while (i <= 10 & x[i] >= 5) {
5 +   print(x[i])
6 +   i = i + 1
7 + }
8 [1] 10
9 [1] 9
10 [1] 8
11 [1] 7
12 [1] 6
13 [1] 5

```

Příklad 5

Pomocí `while` cyklu zjistěte, kolikrát musíme vydělit hodnotu 10, aby byla její aktuální hodnota menší než 0.001. Hodnotu 10 budeme dělit 2.

Řešení 5

```
1 > # Definujeme vychozi hodnotu x a iterator (i)
2 > x <- 10
3 > i <- 0
4 > # Realizace while cyklu
5 > while (x > 0.001) {
6 +   x = x / 2
7 +   i = i + 1
8 + }
9 > # Prvni hodnota x mensi nez 0.001
10 > print(x)
11 [1] 0.0006103516
12 > # Kolikrat bylo potreba hodnotu 10 vydelit - kolikrat probehl
    cyklus
13 > print(i)
14 [1] 14
```

Příklad 6

Máte vektor `x`, který je definován jako sinusoida od 0 do hodnoty π s krokem 0,1. Následně zkuste pomocí `while` cyklu (a vlastnoručně definovaného iterátoru) najít první největší hodnotu, po které začínají hodnoty prvku `x` klesat (tedy hledáme jakési první lokální maximum).

Řešení 6

```
1 > # Definice vektoru:
2 > x <- sin(seq(from = 0, to = pi, by = 0.1))
3 > print(x, digits = 1)
4 [1] 0.00 0.10 0.20 0.30 0.39 0.48 0.56 0.64 0.72 0.78 0.84
5 [12] 0.89 0.93 0.96 0.99 1.00 1.00 0.99 0.97 0.95 0.91 0.86
6 [23] 0.81 0.75 0.68 0.60 0.52 0.43 0.33 0.24 0.14 0.04
7 >
8 > # Iteracni promenna:
9 > i = 2
10 > # while cyklus kde testujeme zda prvek x stale roste:
11 > while (x[i + 1] > x[i]) {
```

```
12 +     i = i + 1
13 + }
14 > # Na ktere i-te pozici je prvek nejvetsi:
15 > print(i)
16 [1] 17
17 > # Jaka je hodnota nejvyssiho prvku:
18 > print(x[i])
19 [1] 0.9995736
```

Správné by bylo v těchto příkladech vyřešit (ošetřit) různé problematické situace (kdy například iterátor převyší délku x). Toto a jiné techniky si ukážeme v následujících kapitolách.

Kapitola 2

Podmínky

Často nastává situace, kdy potřebujeme realizovat část kódu jen za určité podmínky. V takovém případě pak využijeme právě podmínku. Nejpoužívanější podmínkou je funkce `if` a dále funkce `ifelse()`. Podmínku `if` můžeme dále rozšířit o část `else` a `else if`. Pojdme se podívat na jednotlivou problematiku v následujících kapitolách.

2.1 Podmínka `if`

Podmínka `if` prvně ověří, zda je logický výraz v kulaté závorce za zápisem `if` pravdivý, a pokud je, tak **R** vykoná kód ve složených závorkách. Díky tomuto postupu jsme schopni vykonávat určité sekvence kódu jen za určité situace. Ukážeme si prvně podmínku se vždy splněným výrazem.

```
1 > if (TRUE) {  
2 +   print("Splněna")  
3 + }  
4 [1] "Splněna"
```

Naopak podmínka, jejíž kód se ve složené závorce nikdy nerealizuje, by byla:

```
1 > if (FALSE) {  
2 +   print("nesplněna, tento kód se Vám nikdy nezobrazí")  
3 + }
```

Samozřejmě obvykle nebudeme přímo psát výsledek logické operace do kulatých závorek ručně, ale stanovíme logický výraz, který bude testovat danou situaci. Například když chceme, aby se pro čísla od jedné do dvanácti, která jsou celočíselně dělitelná třemi, zobrazil specifický text, můžeme stanovit podmínku pomocí nulového zbytku po dělení.

```

1 > for (i in 1:12) {
2 +   if (i %% 3 == 0) {
3 +     print(paste("Cislo", i, "je celociselne delitene trema"))
4 +   }
5 + }
6 [1] "Cislo 3 je celociselne delitene trema"
7 [1] "Cislo 6 je celociselne delitene trema"
8 [1] "Cislo 9 je celociselne delitene trema"
9 [1] "Cislo 12 je celociselne delitene trema"

```

Výraz `i%%3 == 0` testuje, zda zbytek po dělení řídicí proměnné `for` cyklu je nulový. V případě, kdy je, se pak realizuje část kódu příslušící k tělu podmínky. V našem případě se zobrazí hodnota čísla, které daný test splnilo, a vypíše se informace o tom, že daná hodnota je celočíselně dělitelná třemi.

Samozřejmě podmínky mohou být i komplexnější a můžeme testovat celou sérii logických výrazů. Například můžeme pro čísla jedna až sto otestovat, zda je dané číslo celočíselně dělitelné jak třemi, tak pěti, a pokud je liché, tak jej zobrazíme v konzoli. Řešení pak najdeme pomocí spojení několika logických výrazů do jedné podmínky, jak ilustruje následující kód.

```

1 > for (i in 1:100) {
2 +   if (i%%3 == 0 & i%%5 == 0 & i%%2 != 0) {
3 +     print(i)
4 +   }
5 + }
6 [1] 15
7 [1] 45
8 [1] 75

```

Na tomto příkladu je vidět, že pomocí logických operátorů, jako jsou `&` (and) a `|` (or), můžeme zapsat opravdu komplexní testovací úlohy.

Příklad 1

Pro sekvenci hodnot definovanou jako `seq(from = 1, to = 100, by = 3)` zobrazte hodnoty celočíselně dělitelné pěti ve `for` cyklu.

Řešení 1

```

1 > #Definice promenne
2 > x <- seq(from = 1, to = 100, by = 3)
3 > #Vlastni cyklus
4 > for (i in 1:length(x)) {

```

```

5 + #IF podmínka
6 + if (x[i] %% 5 == 0)
7 +   print(x[i])
8 + }
9 [1] 10
10 [1] 25
11 [1] 40
12 [1] 55
13 [1] 70
14 [1] 85
15 [1] 100

```

2.2 Větvení podmínek pomocí else

Často nastává situace, kdy v případě, že není podmínka splněná, se má realizovat jiný kód, než když podmínka splněná je. Tuto situaci můžeme vyřešit pomocí rozšíření `if` podmínky o výraz `else`. Jedná se tedy o alternativní část kódu, kterou umístíme přímo za tělo podmínky `if`, tedy přímo za druhou složenou závorku. Následující kód bude hodnoty od jedné do desíti testovat na to, zda jsou sudé nebo liché. Tento problém by šel jistě vyřešit i dvěma podmínkami `if`, kdy by jedna testovala, zda je číslo liché, a druhá, zda je sudé. Jedná se ale o velmi špatný zvyk a v případě komplexnějších úloh by byl kód nejen hůře čitelný a delší, ale i by se v něm snáze vyrobila chyba.

```

1 > for (i in 1:10) {
2 +   if (i %% 2 == 0) {
3 +     print(paste("hodnota: ", i, "je suda"))
4 +   } else{
5 +     print(paste("hodnota: ", i, "je licha"))
6 +   }
7 + }
8 [1] "hodnota: 1 je licha"
9 [1] "hodnota: 2 je suda"
10 [1] "hodnota: 3 je licha"
11 [1] "hodnota: 4 je suda"
12 [1] "hodnota: 5 je licha"
13 [1] "hodnota: 6 je suda"
14 [1] "hodnota: 7 je licha"
15 [1] "hodnota: 8 je suda"
16 [1] "hodnota: 9 je licha"
17 [1] "hodnota: 10 je suda"

```

V případě nesplnění podmínky se realizuje kód uvedený v `else` větvi podmínky. Naopak v případě platnosti podmínky již nedojde k realizaci tohoto kódu, ale realizuje se pouze část první (přímo za `if` podmínkou).

Příklad 2

Pomocí `if...else` v datasetu `mtcars` vypište do konzole, zda auto má či nemá automatickou převodovku (proměnná `am`). Daná úloha má být provedena pouze pro auta s šesti válci (proměnná `cyl`).

Řešení 2

```
1 > for (i in 1:nrow(mtcars)) {
2 +   if (mtcars$cyl[i] == 6) {
3 +     if (mtcars$am[i] == 1) {
4 +       print(paste("auto", rownames(mtcars)[i], "ma automatickou
      převodovku"))
5 +     } else{
6 +       print(paste("auto", rownames(mtcars)[i], "nema automatickou
      převodovku"))
7 +     }
8 +   }
9 + }
```

```
10 [1] "auto Mazda RX4 ma automatickou převodovku"
11 [1] "auto Mazda RX4 Wag ma automatickou převodovku"
12 [1] "auto Hornet 4 Drive nema automatickou převodovku"
13 [1] "auto Valiant nema automatickou převodovku"
14 [1] "auto Merc 280 nema automatickou převodovku"
15 [1] "auto Merc 280C nema automatickou převodovku"
16 [1] "auto Ferrari Dino ma automatickou převodovku"
```

2.3 Větvení podmínek pomocí `else if`

Další možné rozšíření `if` podmínky je `else if`. Tento výraz nám umožní rozšířit podmínku `if` tak, aby v případě, kdy nebude první podmínka splněna (v `if` části), byla testována alternativní podmínka (či podmínky v `else if` části). Tedy po nesplnění první podmínky dochází k testování další podmínky, ale v `else if` části. Pokud ani tato podmínka není splněná, tak se přechází k `else` části, anebo je testování zcela ukončeno. Je možné vytvořit podmínku, která bude mít po `if` části více `else if` výrazů. Podstatné je, aby měla pouze jen jednu `else` větev, která je až poslední (po `else if` výrazech). Při testování podmínek hraje roli i pořadí. Podmínky jsou testované v pořadí, v jakém jsou zapsané. Tedy v podmínce může být splněno více podmínek jednotlivých větví kódu, ale realizuje se pouze ta, která je splněna jako první! Ostatní splněné podmínky se vůbec nerealizují.

```

1 > for (i in 1:10) {
2 +   if (i %in% 1:5) {
3 +     print(paste(i, ": a"))
4 +   } else if (i %in% 1:7) {
5 +     print(paste(i, ": b"))
6 +   } else if (i %in% 1:10) {
7 +     print(paste(i, ": c"))
8 +   } else{
9 +     print("ostatni")
10 +   }
11 + }
12 [1] "1 : a"
13 [1] "2 : a"
14 [1] "3 : a"
15 [1] "4 : a"
16 [1] "5 : a"
17 [1] "6 : b"
18 [1] "7 : b"
19 [1] "8 : c"
20 [1] "9 : c"
21 [1] "10 : c"

```

Pro čísla od jedné do desíti testujeme následující podmínku. V případě, kdy je číslo přítomné v množině čísel jedna až pět, pak zobrazíme písmeno „a“. Pokud je číslo přítomné v množině čísel jedna až sedm, tak zobrazíme písmeno „b“, a pokud je číslo přítomné v množině jedna až deset, tak zobrazíme písmeno „c“, jinak zobrazíme text „ostatní“. Všimněte si, že například hodnota pět splňuje všechny podmínky, ale zobrazí se u ní pouze písmeno „a“. Toto je dané právě pořadím podmínek.

Jak byste pomocí předchozích znalostí otestovali čísla jedna až deset na to, zda jsou prvočísla? V **R** to lze udělat jednoduše.

```

1 > for (i in 1:10) {
2 +   if (i == 2) {
3 +     print(paste("Cislo: ", i, "je prvocislo"))
4 +   } else if (any(i %% c(2:(i-1)) == 0)) {
5 +     print(paste("Cislo: ", i, "neni prvocislo"))
6 +   } else {
7 +     print(paste("Cislo: ", i, "je prvocislo"))
8 +   }
9 + }
10 [1] "Cislo: 1 neni prvocislo"
11 [1] "Cislo: 2 je prvocislo"
12 [1] "Cislo: 3 je prvocislo"
13 [1] "Cislo: 4 neni prvocislo"
14 [1] "Cislo: 5 je prvocislo"

```

```
15 [1] "Cislo: 6 není prvocislo"
16 [1] "Cislo: 7 je prvocislo"
17 [1] "Cislo: 8 není prvocislo"
18 [1] "Cislo: 9 není prvocislo"
19 [1] "Cislo: 10 není prvocislo"
```

Cyklus prvně otestuje, zda je číslo rovno dvěma, pokud je, tak číslo je prvočíslem. Pokud není, tak se testuje, zda číslo má jiného dělitele než jedna a sebe samého (tedy od dvou do $i - 1$). Pokud má (je tedy tímto číslem celočíselně dělitelné – zbytek po dělení je roven nule), tak je číslo prohlášené za hodnotu, která není prvočíslem. Všechna ostatní čísla (else část) jsou prvočísla (jako číslo dva).

Příklad 3

Zkuste předchozí příklad (č. 3) vyřešit pomocí jen jedné podmínky s využitím `else if` větve podmínky.

Řešení 3

```
1 > for (i in 1:nrow(mtcars)) {
2 +   # Do if casti je nutne pridat dalsi podminku na pocet valcu
3 +   if (mtcars$am[i] == 1 & mtcars$cyl[i] == 6) {
4 +     print(paste("auto", rownames(mtcars)[i], "ma automatickou
      prevodovku"))
5 +     # a dale je nutne i v else if casti osetrit podminku na pocet
      valcu
6 +   } else if (mtcars$cyl[i] == 6) {
7 +     print(paste("auto", rownames(mtcars)[i], "nema automatickou
      prevodovku"))
8 +   }
9 + }
10 [1] "auto Mazda RX4 ma automatickou prevodovku"
11 [1] "auto Mazda RX4 Wag ma automatickou prevodovku"
12 [1] "auto Hornet 4 Drive nema automatickou prevodovku"
13 [1] "auto Valiant nema automatickou prevodovku"
14 [1] "auto Merc 280 nema automatickou prevodovku"
15 [1] "auto Merc 280C nema automatickou prevodovku"
16 [1] "auto Ferrari Dino ma automatickou prevodovku"
```

2.4 Větvení podmínek ifelse

V případě, kdy je kód, který se má realizovat v podmínce, velmi jednoduchý, lze alternativně využít funkci `ifelse()`. Funkce má tyto parametry: `test`, `yes`, `no`. Tedy výraz, který testuje (`test`), `yes` – co se má stát v případě, kdy je podmínka splněna, `no` – co se má stát v případě, kdy podmínka splněna není.

```
1 > for (i in 1:6) {  
2 +   ifelse(i <= 3, print("mensi nebo rovno trem"), print("vetsi nez  
   3"))  
3 + }  
4 [1] "mensi nebo rovno trem"  
5 [1] "mensi nebo rovno trem"  
6 [1] "mensi nebo rovno trem"  
7 [1] "vetsi nez 3"  
8 [1] "vetsi nez 3"  
9 [1] "vetsi nez 3"
```

Zajímavostí je, že lze v rámci této funkce deklarovat (přiřazovat) hodnoty i do proměnných.

```
1 > ifelse(TRUE, x <- 10, x <- 100)  
2 [1] 10  
3 > x  
4 [1] 10
```

Nejedná se ale o doporučený postup. V případě deklarací proměnných či komplexnějších kódů je lepší zvolit standardní `if...else` výraz.

Příklad 4

Pomocí `for` cyklu a `ifelse` funkce nahraďte v sekvenci `sin(seq(from = 0, to = pi*4, by = 0.5))` hodnoty větší než nula hodnotou jedna a menší nebo rovno nule hodnotou nula.

Řešení 4

```
1 > # Deklarace promenne  
2 > x <- sin(seq(from = 0, to = pi * 4, by = 0.5))  
3 > print(x, digit = 1)  
4 [1] 0.00 0.48 0.84 1.00 0.91 0.60 0.14 -0.35 -0.76 -0.98  
   -0.96 -0.71 -0.28 0.22  
5 [15] 0.66 0.94 0.99 0.80 0.41 -0.08 -0.54 -0.88 -1.00 -0.88  
   -0.54 -0.07
```

```
6 > # for cyklus
7 > for(i in 1:length(x)) ifelse(x[i] > 0, x[i] <- 1, x[i] <- 0)
8 > print(x)
9 [1] 0 1 1 1 1 1 1 0 0 0 0 0 0 1 1 1 1 1 0 0 0 0 0 0
```

Kapitola 3

Příkazy next, break a repeat cyklus

V této kapitole se seznámíme s výrazy `next`, `break` a `repeat` cyklem. `Next` a `break` jsou takzvané řídicí výrazy a umožňují nám ovládat cyklus mimo předem stanovený běh cyklu. Dále se seznámíme s cyklem `repeat`. `Repeat` cyklus je takzvaný nekonečný cyklus. Pokud nestanovíme ukončení pomocí výrazu `break`, tak tento cyklus poběží „do nekonečna“. Z tohoto důvodu je daná problematika spojena do jedné kapitoly.

3.1 Příkaz next a break

Výrazy `next` a `break` nám umožňují řídit cyklus. Výraz `next` způsobí to, že daná iterace cyklu se od místa zápisu `next` ukončí a začne ihned iterace následující. Naopak výraz `break` ukončí rovnou celý cyklus jako takový. Oba dva výrazy lze kombinovat a uplatnit je v jakémkoliv cyklu. Tedy lze je uplatnit jak ve `for`, tak `while` i `repeat` cyklu.

Představme si, že máme data od jedné do sta a chceme zastavit cyklus na místě, kdy součet všech předchozích hodnot je více než padesát.

```
1 > for (i in 1:100) {  
2 +   if (sum(1:i) > 50) break  
3 +   print(paste("iterace:", i, "suma:", sum(1:i)))  
4 + }  
5 [1] "iterace: 1 suma: 1"  
6 [1] "iterace: 2 suma: 3"  
7 [1] "iterace: 3 suma: 6"  
8 [1] "iterace: 4 suma: 10"  
9 [1] "iterace: 5 suma: 15"  
10 [1] "iterace: 6 suma: 21"
```

```
11 [1] "iterace: 7 suma: 28"
12 [1] "iterace: 8 suma: 36"
13 [1] "iterace: 9 suma: 45"
```

Množina for cyklu, přes kterou má cyklus iterovat, je od jedné do sta, avšak cyklus se zastaví již v desáté iteraci. V deváté iteraci byl součet ještě čtyřicet pět, a proto cyklus pokračoval dále, ale již v desáté iteraci byl větší než padesát, a tak se cyklus od místa podmínky ukončil. Tedy sice do desáté iterace vstoupil, ale již ji nevykonal celou, ještě před funkcí `print` se zastavil. Oba výrazy můžeme i kombinovat. Například přidat podmínku pro případ, že se daná iterace přeskočí, když součet bude lichý.

```
1 > for (i in 1:100) {
2 +   if (sum(1:i) > 50) {
3 +     break
4 +   } else if (sum(1:i) %% 2 != 0) {
5 +     next
6 +   }
7 +   print(paste("iterace:", i, "suma:", sum(1:i)))
8 + }
9 [1] "iterace: 3 suma: 6"
10 [1] "iterace: 4 suma: 10"
11 [1] "iterace: 7 suma: 28"
12 [1] "iterace: 8 suma: 36"
```

Jak z výsledků vidíme, hodnoty 1, 3, 15, 21, a 45 cyklus při informaci o tom, jaké hodnoty splnily podmínky, vynechal. Cyklus pokračoval dále, ale od místa testu `else if` daná iterace vždy skončila. Zkusme se teď podívat na problematický `while` cyklus, který má špatně nastavenou podmínku, a tak nikdy nekončí. Někdy se při implementaci matematických metod může stát, že si nejsme zcela jisti, že může nastat situace, kdy je podmínka nesplněna. Pak se hodí použití `break` řídicího výrazu v kombinaci s vlastním iterátorem. Následující kód by doběhl po velmi dlouhé době, ale díky zapojení iterátoru jsme se na tuto situaci připravili, a v případě, kdy se nám nepodaří při určitém počtu iterací dosáhnout negativní podmínky, tak cyklus ukončí výraz `break`.

```
1 > iterator = 1
2 > x = 0
3 > while (abs(x) < 10) {
4 +   x + 1e-5
5 +   iterator = iterator + 1
6 +   if (iterator >= 100) break
7 + }
8 >
9 > print(iterator)
10 [1] 100
```

3.2 Cyklus repeat

Cykly `for` a `while` nejsou jedinými cykly, které jsou v jazyku **R** dostupné. Alternativou k těmto cyklům je použití obecnějšího cyklu `repeat`. Cyklus `for` má přesně definovaný počet iterací, kolikrát se vykoná. Tuto informaci má cyklus z délky množiny, kterou přebírá jako vstupní parametr. Cyklus `while` naopak před začátkem každého cyklu ověřuje platnost podmínky, pokud není splněna, tak se cyklus ukončí. Tedy oba cykly mají definován způsob ukončení cyklu. Cyklus `repeat` nemá předem definován způsob ukončení cyklu a ani jinak není omezen ve svém opakování. Jedná se záměrně o nekonečnou smyčku. Cyklus lze připodobnit `while` cyklu s nastavenou podmínkou na `TRUE`, tedy `while(TRUE)`. Uživatel pak musí sám definovat způsob ukončení cyklu v samotném těle cyklu. Ukončení cyklu se dosáhne pomocí podmínky `if` v kombinaci s řídicím výrazem `break`. Zkusme si tedy zapsat variantu `for` cyklu pomocí `break` cyklu. Následující kód ilustruje použití `repeat` cyklu ve smyslu `for` cyklu. Máme zde řídicí proměnnou (`i`), kterou iterujeme přes množinu jedna až deset. Navýšení hodnoty proměnné je prováděno o hodnotu jedna a poslední hodnotou je hodnota deset (definovaná proměnnou `pocet_cyklu`).

```
1 > i <- 0
2 > pocet_cyklu = 10
3 > repeat {
4 +   print(i)
5 +   if (i >= pocet_cyklu) {
6 +     break
7 +   } else {
8 +     i = i + 1
9 +   }
10 + }
11 [1] 0
12 [1] 1
13 [1] 2
14 [1] 3
15 [1] 4
16 [1] 5
17 [1] 6
18 [1] 7
19 [1] 8
20 [1] 9
21 [1] 10
```

Alternativně si lze představit i variantu `while` cyklu. Zde provádíme výpočet: $x = (x + \pi) * x / \pi$, a pokud hodnota x převyší hodnotu 1000, tak se cyklus zastaví.

```
1 > x <- 1
2 > repeat {
3 +   x = (x + pi) * x / pi
```



```

4 +   print(x)
5 +   if (x > 1000)
6 +       break
7 + }
8 [1] 1.31831
9 [1] 1.871514
10 [1] 2.986414
11 [1] 5.825315
12 [1] 16.62694
13 [1] 104.6253
14 [1] 3588.99

```

Zde je důležité si uvědomit, že hodnota `x` je v posledním kole cyklu vyšší než stanovená podmínka. Tato situace je dána tím, že testujeme podmínku `if` s příkazem `break` až po výpočtu nové hodnoty `x`.

Příklad 1

Pomocí cyklu `repeat` spočítejte variační koeficient pro všechny numerické proměnné datasetu `iris`.

Řešení 1

```

1 > i = 1
2 > repeat {
3 +   # zastavení cyklu
4 +   if (i > ncol(iris)) break
5 +   # samotné tělo cyklu
6 +   if (is.numeric(iris[, i])) {
7 +       print(sd(iris[, i]) / mean(iris[, i]) , digit = 3)
8 +   }
9 +   # navýšení hodnoty iteratoru
10 +   i = i + 1
11 + }
12 [1] 0.142
13 [1] 0.143
14 [1] 0.47
15 [1] 0.636

```

Kapitola 4

Funkce `apply`, `by` a `aggregate`

Rodina funkcí `apply` slouží k automatizaci určité funkce pro určitý objekt, popřípadě jeho část (např. část dimenze u matice – sloupce, řádky...). `Apply` funkce nepatří u začátečníků mezi oblíbené funkce, ale pokud se je naučíme používat, zkrátí nám nejen čas věnovaný zápisu kódu, ale také zjednoduší kód a mnohdy ho zrychlí. Jedná se o následující funkce:

- `apply()`,
- `sapply()`,
- `vapply()`,
- `lapply()`,
- `mapply()`,
- `rapply()`,
- `tapply()`.

K `apply` funkcím se standardně řadí navíc i funkce `by()` a `aggregate()`. Tyto funkce mají trochu odlišné použití, a ukážeme si to na konci této kapitoly.

4.1 Funkce `apply`

Základní funkcí, kterou musíte poznat, je funkce `apply`. Všechny ostatní navazující funkce jsou totiž její variantou. Funkce `apply()` má tři povinné parametry. Jedná se o: `X`, `MARGIN` a `FUN`.

Do parametru **X** předáváme objekt, pro který se má funkce provést. Parametr **MARGIN** definuje dimenzi, skrze kterou se má funkce provést. Poslední povinný parametr je **FUN**, do kterého předáváme jméno funkce, která se má přes specifikovanou dimenzi daného objektu provést.

Parametr označující dimenzi (**MARGIN**) je vhodné blíže osvětlit. Vrátime-li se na chvíli k výběru elementů jednotlivých objektů, například vektor má jednu dimenzi, matice dvě a pole může mít i tři a více dimenzí (ale i jednu nebo dvě). Když provádíme výběr prvků pomocí hranatých závorek, tak ve vektoru **v** (například o deseti prvcích) vybereme druhý pomocí příkazu **v[2]**. V matici (**V**, například o deseti řádkách a deseti sloupcích) bychom výběr prvku v druhém řádku a druhém sloupci provedli pomocí výrazu **V[2,2]**. V případě trojrozměrného pole (pro zjednodušení příkladu si teď uveďme pole o rozměrech $3 \times 3 \times 3$), kdybychom chtěli vybrat prvek v druhém řádku, druhém sloupci a druhé řadě, tak to provedeme pomocí výrazu: **R[2,2,2]**. V případě vektoru **v** se tedy uvádí jen pořadí prvku, jenž chceme vybrat (**v[číslo prvku]**), v případě matice **V** jsme zapsali řádku, sloupec (**V[číslo řádky, číslo sloupce]**) a v případě našeho pole jsme zapsali všechny tři dimenze (**R[číslo řádky, číslo sloupce, číslo řady]**). Do parametru funkce **apply** dimenze (**MARGIN**) pak zapisujeme pořadí dané dimenze při výběru. V případě aplikace funkce **apply** na matici, kdy budou hodnoty parametru **MARGIN** 1, tak se daná funkce provede přes řádky daného objektu (řádky se zapisují jako první do hranatých závorek). Naopak, kdybychom zvolili hodnotu dva, tak by se daná funkce aplikovala přes sloupce daného objektu, protože výběr přes sloupce je až druhá hodnota, co se zapisuje do hranatých závorek. S podobnou úvahou funguje tato funkce i pro pole a jiné objekty.

S ohledem na to, že se jedná o dost abstraktní definici toho, jak funkce funguje, uveďme si jednoduchý příklad. Když bychom chtěli spočítat průměr pro každý sloupec matice, tak máme v jazyku **R** několik možností:

- Použijeme funkci **colMeans()**.
- Vytvoříme cyklus, který si postupně vybere sloupec po sloupci dané matice a spočítá z něj průměr.
- Použijeme funkci **apply**, které předáme parametr hrany (dimenze, přes kterou sčítáme), funkci, kterou použijeme (**mean()**), a objekt, na který chceme funkci aplikovat.

Následující kód ilustruje řešení úlohy výpočtem průměru dat přes sloupce/řádky pomocí výchozích funkcí jazyka **R** **colMeans()**/**rowMeans()**.

```
1 > V <- matrix(data = c(1:9),
2 +             nrow = 3,
3 +             ncol = 3)
4
5 > print(V)
```

```

6      [,1] [,2] [,3]
7 [1,]    1    4    7
8 [2,]    2    5    8
9 [3,]    3    6    9
10
11 > prumer_sloupcu <- colMeans(V)
12
13 > prumer_radku <- rowMeans(V)
14
15 > print(prumer_sloupcu)
16 [1] 2 5 8
17
18 > print(prumer_radku)
19 [1] 4 5 6

```

Alternativní řešení lze samozřejmě provést pomocí for cyklu. Zde je důležité poznamenat, že se jedná o dost neefektivní řešení. V případě, kdy neexistuje výchozí funkce v jazyku **R** (například pro výpočty intervalů spolehlivosti pro celý dataset), tak mnoho uživatelů volí tento postup. Efektivnějším řešením by ale bylo použít `apply()` funkci.

```

1 > V <- matrix(data = c(1:9),
2 +           nrow = 3,
3 +           ncol = 3)
4
5 > print(V)
6      [,1] [,2] [,3]
7 [1,]    1    4    7
8 [2,]    2    5    8
9 [3,]    3    6    9
10
11 > prumer_sloupcu <- rep(NA, ncol(V))
12
13 > prumer_radku <- rep(NA, nrow(V))
14
15 > for (i in 1:ncol(V)) prumer_sloupcu[i] <- mean(V[, i])
16
17 > for (i in 1:nrow(V)) prumer_radku[i] <- mean(V[i, ])
18
19 > print(prumer_sloupcu)
20 [1] 2 5 8
21
22 > print(prumer_radku)
23 [1] 4 5 6

```

Řešení úlohy pomocí `apply` funkce pak představuje následující kód. Rozdíl ve výpočtu průměru přes řádek/sloupec pak spočívá v čísle hrany, přes kterou se daná funkce realizuje.

```

1 > V <- matrix(data = c(1:9),
2 +             nrow = 3,
3 +             ncol = 3)
4 > print(V)
5      [,1] [,2] [,3]
6 [1,]    1    4    7
7 [2,]    2    5    8
8 [3,]    3    6    9
9 > prumer_sloupcu <- apply(X = V, MARGIN = 2, FUN = mean)
10 > prumer_radku <- apply(X = V, MARGIN = 1, FUN = mean)
11 > print(prumer_sloupcu)
12 [1] 2 5 8
13 > print(prumer_radku)
14 [1] 4 5 6

```

Jak jsme si již představili, v případě průměru má **R** výchozí funkci, která danou úlohu provede efektivněji (`colMeans()`/`rowMeans()`). Jazyk **R** má dále obdobnou funkci pro sumu (`rowSums()`/`colSums()`). Existuje ale mnoho dalších funkcí, které **R** nemá převedené pro maticové/tabulkové výpočty – přes všechny sloupce či řádky. Ukažme si například způsob, jak vypočítat medián přes všechny proměnné datasetu `mtcars`.

```

1 > apply(X = mtcars, MARGIN = 2, FUN = median)
2      mpg      cyl      disp      hp      drat      wt      qsec      vs
3 19.200    6.000 196.300 123.000    3.695    3.325 17.710    0.000
   0.000    4.000    2.000

```

Obdobně můžeme spočítat i rozptyl či směrodatnou odchylku:

```

1 > apply(X = mtcars, MARGIN = 2, FUN = sd)
2      mpg      cyl      disp      hp
3 6.0269481 1.7859216 123.9386938 68.5628685
4      drat      wt      qsec      vs
5 0.5346787 0.9784574 1.7869432 0.5040161
6      am      gear      carb
7 0.4989909 0.7378041 1.6152000

```

`Apply` tedy postupně dosazuje za první parametr funkce, definované v parametru `FUN`, výběr dat specifikovaný dimenzí. `Apply` umožňuje specifikovat uživateli ostatní parametry funkce, do které dosazuje. Tyto parametry lze zapsat na konec funkce `apply`, je ale nutné je specifikovat správným jménem. Například si vytvoříme matici, která bude obsahovat `NA` řádek. Bez specifikace parametru `na.rm` ve funkci `mean` **R** nevrátí požadovaný průměr, ale `NA` hodnotu.

```

1 > A <- matrix(1:9, 3)

```

```

2 > A[2, ] <- NA
3 > apply(X = A, MARGIN = 2, FUN = mean)
4 [1] NA NA NA
5 > apply(X = A, MARGIN = 2, FUN = mean, na.rm = TRUE)
6 [1] 2 5 8

```

Tak jako v případě matice či data frame objektu lze aplikovat funkci `apply` i na objekt typu pole (`array`). Vytvoříme si pole velikosti $3 \times 3 \times 3$, tedy pole o třech dimenzích, a zkusíme si provést jednotlivé možné součty. Prvně si zkusíme vytvořit řádkový součet skrze celé pole. Zůstane zachována dimenze řádku a všechny ostatní dimenze budou sečteny. Toto provedeme pomocí parametru `MARGIN = 1`.

```

1 > pole = array(data = c(1:27), dim = c(3, 3, 3))
2 > apply(X = pole, MARGIN = 1, FUN = sum)
3 [1] 117 126 135

```

V případě pole nemusí být aplikace funkce `apply` na první pohled srozumitelná, proto si ukážeme i řešení dané úlohy pomocí funkcí `rowSums()` a selekcí příslušných dimenzí pole.

```

1 > pole = array(data = c(1:27), dim = c(3, 3, 3))
2 > cbind(sum(pole[1,,]), sum(pole[2,,]), sum(pole[3,,]))
3 [1] 117 126 135

```

Co kdybychom ale chtěli vytvořit matici součtů řádků skrze všechny dimenze a ne součet řádkových sum? V tomto případě je nutné do parametru `MARGIN` specifikovat všechny dimenze, které mají být zachovány. V našem případě půjde o dimenzi, skrze kterou provádíme součet a dále zbývající (třetí) dimenzi úlohy, aby zůstala zachována všechna data.

```

1 > pole = array(data = c(1:27), dim = c(3, 3, 3))
2 > apply(X = pole, MARGIN = c(1, 3), FUN = sum)
3      [,1] [,2] [,3]
4 [1,]   12   39   66
5 [2,]   15   42   69
6 [3,]   18   45   72

```

Jaké by bylo alternativní řešení?

```

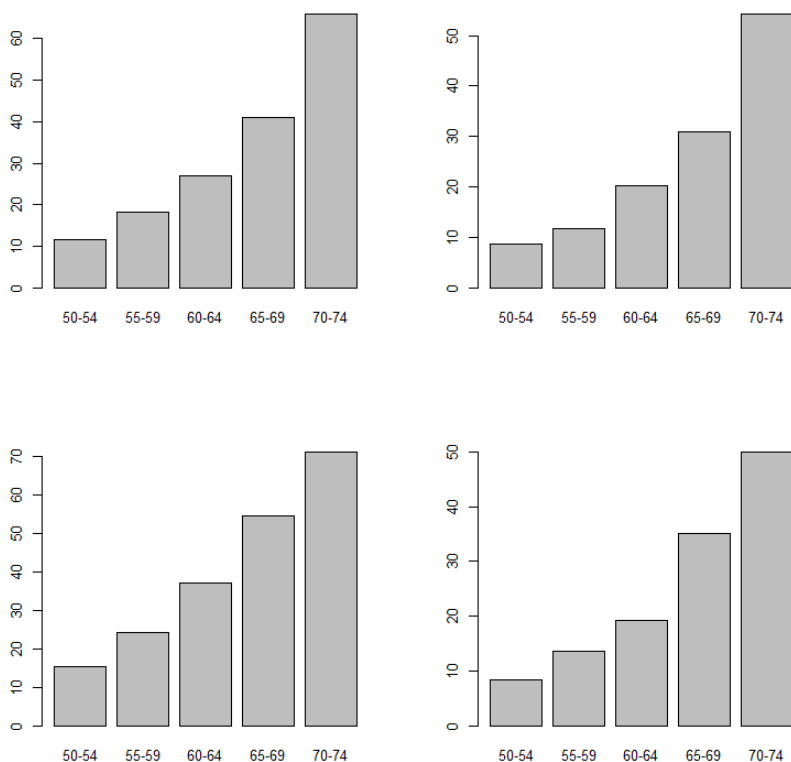
1 > pole = array(data = c(1:27), dim = c(3, 3, 3))
2 > cbind(rowSums(pole[, 1]), rowSums(pole[, 2]), rowSums(pole[, 3]))
3      [,1] [,2] [,3]
4 [1,]   12   39   66
5 [2,]   15   42   69
6 [3,]   18   45   72

```

Zajímavostí je, že lze v **R** zpracovávat pomocí `apply` funkcí i grafické funkce, například vykreslit sérii grafů. Jen je důležité si uvědomit, že `apply()` funkce navrácí návratovou

hodnotu grafické funkce do konzole. Obvyklá návratová hodnota je v případě grafů v jazyku **R** `NULL` nebo vstupní data. Hodnota se nám pak vypíše v případě grafů do příkazové řádky. Tento postup si ilustrujme na datech o počtu úmrtí na 1000 obyvatel ve Virginii v USA. Data jsou tříděná dle věku a místa, kde daný jedinec bydlí.

```
1 > print(VADeaths)
2      Rural Male Rural Female Urban Male Urban Female
3 50-54      11.7      8.7      15.4      8.4
4 55-59      18.1      11.7      24.3      13.6
5 60-64      26.9      20.3      37.0      19.3
6 65-69      41.0      30.9      54.6      35.1
7 70-74      66.0      54.3      71.1      50.0
8 > par(mfrow = c(2, 2))
9 > apply(X = VADeaths, 2, FUN = barplot)
10      Rural Male Rural Female Urban Male Urban Female
11 [1,]      0.7      0.7      0.7      0.7
12 [2,]      1.9      1.9      1.9      1.9
13 [3,]      3.1      3.1      3.1      3.1
14 [4,]      4.3      4.3      4.3      4.3
15 [5,]      5.5      5.5      5.5      5.5
16 > par(mfrow=c(1, 1))
```



Druhé zobrazení dat v příkazové řádce je pak důsledkem aplikace funkce `barplot` ve funkci `apply`. Kdybychom chtěli specifikovat dodatečné parametry grafické funkce `barplot()`, bylo by vhodné použít jinou funkci rodiny `apply`, například `mapply()`, kterou si představíme v navazující podkapitole.

Do všech funkcí rodiny `apply` je možné zadávat i anonymní funkce. Tedy funkce, které nemají jméno a uživatel je zadefinuje přímo při volání `apply` funkce. Více o anonymních funkcích si povíme v kapitole uživatelské funkce. Ukažme si ale dva příklady.

```

1 > apply(X = mtcars, MARGIN = 2, FUN = function(z) max(z) - min(z))
2   mpg      cyl    disp    hp  drat    wt    qsec    vs
3 23.500  4.000 400.900 283.000  2.170  3.911  8.400  1.000
4   gear      carb
5  2.000  7.000

```

V následujícím příkladu je vypočtena dolní a horní mez intervalu spolehlivosti pro data z datasetu `mtcars`. Vyjděme z předpokladu, že data mají normální rozdělení bez známého

rozptylu. Zajímavé na této úloze je to, že funkce `apply()` zvládne navrátit i více hodnot pro daný sloupec, ale musí mít stejnou délku. Dále si všimněte, jak byl omezen výběr dat (pomocí selekce sloupců) přímo v těle samotné funkce.

```
1 > apply(X = mtcars[, c(1, 3:7)], MARGIN = 2,
2 +     FUN = function(x)
3 +         mean(x) + c(-1, +1) *
4 +         (qt(p = 0.975, df = (length(x) - 1)) / sqrt(length(x))))
5 + )
6           mpg      disp      hp      drat      wt      qsec
7 [1,] 19.73009 230.3613 146.327 3.236024 2.856712 17.48821
8 [2,] 20.45116 231.0824 147.048 3.957101 3.577788 18.20929
```

Více si anonymní funkce rozebereme ale až v následujících kapitolách.

Příklad 1

V datasetu `mtcars` spočítejte pro každou proměnnou rozdíl mezi maximální a minimální hodnotou (tedy variační rozpětí).

Řešení 1

```
1 > apply(mtcars, 2, FUN = function(x) max(x) - min(x))
2           mpg      cyl      disp      hp      drat      wt      qsec      vs
3           am      gear      carb
23.500    4.000 400.900 283.000    2.170    3.911    8.400    1.000
1.000     2.000    7.000
```

Příklad 2

Některé datové objekty jsou občas spravovány jako vícerozměrná pole (`array`). V následujícím příkladu si ukážeme, jak pro třídimenzionální pole provést základní operaci – sečtení dimenzí.

Máte třídimenzionální krychli definovanou jako: `krychle = array(data = 1:27, dim = c(3, 3, 3))`. Jak byste provedli součet skrze dimenze tak, aby byly „nasečteny“ dimenze skrze jednotlivé řady? Tedy aby výsledkem byl součet skrze třetí dimenzi. Tedy chceme pomocí `apply()` funkce a funkce `sum()` získat výstup analogický k následujícímu kódu:

```
1 > krychle[, , 1] + krychle[, , 2] + krychle[, , 3]
2           [,1] [,2] [,3]
3 [1,]    30    39    48
4 [2,]    33    42    51
5 [3,]    36    45    54
```

Řešení 2

Podíváme se nejprve na to, jak naše data vypadají.

```
1 > krychle = array(data = 1:27, dim = c(3, 3, 3))
2 > print(krychle)
3 , , 1
4
5      [,1] [,2] [,3]
6 [1,]    1    4    7
7 [2,]    2    5    8
8 [3,]    3    6    9
9
10 , , 2
11
12      [,1] [,2] [,3]
13 [1,]   10   13   16
14 [2,]   11   14   17
15 [3,]   12   15   18
16
17 , , 3
18
19      [,1] [,2] [,3]
20 [1,]   19   22   25
21 [2,]   20   23   26
22 [3,]   21   24   27
```

Prostým zadáním 3 do parametru MARGIN získáme zcela jiný výsledek, než byste očekávali.

```
1 > # Prostým zadáním 3 dimenze se sectou uplne cele dimenze ale ne po
   prvcich:
2 > apply(krychle, MARGIN = 3, FUN = sum)
3 [1]  45 126 207
4 > Toto je analogicke k:
5 > c(sum(krychle[, , 1]), sum(krychle[, , 2]), sum(krychle[, , 3]))
6 [1]  45 126 207
```

Řešení je vložení funkce `apply()` do samotné funkce `apply()` tak, jak ukazuje následující kód.

```
1 > apply(krychle, MARGIN = 2, FUN = apply, 1, sum)
2      [,1] [,2] [,3]
3 [1,]   30   39   48
4 [2,]   33   42   51
5 [3,]   36   45   54
```

4.2 Funkce lapply, sapply, vapply, replicate

4.2.1 Funkce lapply

Funkce `lapply()` oproti `apply()` nepracuje již s maticí/polem, ale naopak s vektorem dat. Může se jednat o atomický vektor nebo o vektor typu list (protože i list je vektorem, ale ne atomickým). Dále funkce `lapply()` pracuje s jakýmkoliv objektem, jenž lze převést na list pomocí funkce `as.list()`. Funkce `lapply()` má parametry: `X` – objekt, pro který se provádí, `FUN` – funkce, která se provádí. Použití funkce je pak identické s funkcí `apply()`, ale s tím rozdílem, že se zde již nespecifikuje dimenze dané úlohy (parametr `margin`), funkce je prováděna pro celé prvky listu. Vytvořme si následující list.

```
1 > L<-list(a = c(1:10), b = c(2:5),
2 +       c = c(-5:5), M = matrix(data = c(1:9),ncol = 3))
3 > print(L)
4 $a
5  [1]  1  2  3  4  5  6  7  8  9 10
6
7 $b
8  [1]  2  3  4  5
9
10 $c
11  [1] -5 -4 -3 -2 -1  0  1  2  3  4  5
12
13 $M
14      [,1] [,2] [,3]
15 [1,]    1    4    7
16 [2,]    2    5    8
17 [3,]    3    6    9
```

Funkce `sum()` se použije na všechny objekty zvlášť, protože není nikde zaručeno, že lze převést výsledek funkce na homogenní objekt z pohledu dat i dimenzí, tak je návratový objekt typu list.

Když pak použijeme funkci `lapply()` s funkcí `sum()`, tak funkce `lapply()` spočítá sumu přes celé objekty. Zajímavostí je ale i návratová hodnota. Objekt, který se navrátí, je taktéž typu list.

```
1 > lapply(X = L, FUN = sum)
2 $a
3  [1] 55
4
5 $b
6  [1] 14
7
8 $c
```

```
9 [1] 0
10
11 $M
12 [1] 45
```

Zajímavé na funkci `lapply()` je i to, že ji lze kombinovat s funkcí `apply()` a provést úlohu např. pro všechny proměnné objektu typu `data.frame`. Jediným problémem zde je, že funkce `apply` má stejná jména parametrů jako funkce `lapply()`. Tuto situaci lze vyřešit tak, že parametry pro funkci `apply` se zadají až po parametrech pro funkci `lapply()` a navíc bez jména parametru (nutná podmínka pro kolizní situace parametrů). Musíme proto vypsat parametry ve správném pořadí. Vytvoříme si list matic, pro který budeme chtít spočítat mediánovou hodnotu přes sloupce.

```
1 > L2 <- list(M1 = matrix(1:9, 3),
2 +           M2 = matrix(9:1, 3),
3 +           M3 = matrix(11:19, 3))
4 > print(L2)
5 $M1
6      [,1] [,2] [,3]
7 [1,]    1    4    7
8 [2,]    2    5    8
9 [3,]    3    6    9
10
11 $M2
12      [,1] [,2] [,3]
13 [1,]    9    6    3
14 [2,]    8    5    2
15 [3,]    7    4    1
16
17 $M3
18      [,1] [,2] [,3]
19 [1,]   11   14   17
20 [2,]   12   15   18
21 [3,]   13   16   19
22 > lapply(X = L2, FUN = apply, 2, median)
23 $M1
24 [1] 2 5 8
25
26 $M2
27 [1] 8 5 2
28
29 $M3
30 [1] 12 15 18
```

Ne všem může vyhovovat výstupní formát typu list. V tuto chvíli máme dvě možnosti, jak výstup z funkce `apply` převést např. na matici či pole. První způsob je pomocí funkce

`unlist()`. Funkce „odlistuje“ jakýkoliv list, jednotlivé elementy pojmenuje podle jména listu a podle indexu daného místa na dané pozici v objektu listu.

```
1 > L2 <- list(M1 = matrix(1:9, 3),
2 +           M2 = matrix(9:1, 3),
3 +           M3 = matrix(11:19, 3))
4 > print(L2)
5 $M1
6      [,1] [,2] [,3]
7 [1,]    1    4    7
8 [2,]    2    5    8
9 [3,]    3    6    9
10
11 $M2
12      [,1] [,2] [,3]
13 [1,]    9    6    3
14 [2,]    8    5    2
15 [3,]    7    4    1
16
17 $M3
18      [,1] [,2] [,3]
19 [1,]   11   14   17
20 [2,]   12   15   18
21 [3,]   13   16   19
22 > unlist(lapply(X = L2, FUN = apply, 2, median))
23
24 M11 M12 M13 M21 M22 M23 M31 M32 M33
25   2   5   8   8   5   2  12  15  18
```

Dalším způsobem, jak převést list na jiný objekt, je za použití funkce `simplify2array()`. Pokud je délka jednotlivých elementů listu shodná, pak daný objekt funkce převede na pole.

```
1 > L2 <- list(M1 = matrix(1:9, 3),
2 +           M2 = matrix(9:1, 3),
3 +           M3 = matrix(11:19, 3))
4 > print(L2)
5 $M1
6      [,1] [,2] [,3]
7 [1,]    1    4    7
8 [2,]    2    5    8
9 [3,]    3    6    9
10
11 $M2
12      [,1] [,2] [,3]
13 [1,]    9    6    3
14 [2,]    8    5    2
```

```

15 [3,]      7      4      1
16
17 $M3
18      [,1] [,2] [,3]
19 [1,]    11    14    17
20 [2,]    12    15    18
21 [3,]    13    16    19
22 > simplify2array(lapply(X = L2, FUN = apply, 2, median))
23      M1 M2 M3
24 [1,]   2   8 12
25 [2,]   5   5 15
26 [3,]   8   2 18

```

4.2.2 Funkce sapply

Funkce `sapply()` je formou `lapply()` funkce s tím rozdílem, že výstupní hodnotou není list, ale přímo pole/matice, pokud to lze. Pokud to možné není, tak funkce `sapply()` navrací stejně jako funkce `lapply()` list. Ukážeme si příklad, kdy první list obsahuje tři matice 3×3 a druhý list (L3) obsahuje místo toho jednu matici o velikosti 9×1 .

```

1 > L2 <- list(M1 = matrix(1:9, 3),
2 +           M2 = matrix(9:1, 3),
3 +           M3 = matrix(11:19, 3))
4 > print(L2)
5 $M1
6      [,1] [,2] [,3]
7 [1,]     1     4     7
8 [2,]     2     5     8
9 [3,]     3     6     9
10
11 $M2
12      [,1] [,2] [,3]
13 [1,]     9     6     3
14 [2,]     8     5     2
15 [3,]     7     4     1
16
17 $M3
18      [,1] [,2] [,3]
19 [1,]    11    14    17
20 [2,]    12    15    18
21 [3,]    13    16    19
22 > sapply(X=L2, FUN=sum)
23      M1 M2 M3
24    45  45 135
25

```

```

26 > sapply(X=L2 ,FUN=apply,2, median)
27      M1 M2 M3
28 [1,]  2  8 12
29 [2,]  5  5 15
30 [3,]  8  2 18
31
32 > L3 <- list(M1 = matrix(1:9, 3),
33 +           M2 = matrix(9:1, 9),
34 +           M3 = matrix(11:19, 3))
35 > print(L3)
36 $M1
37      [,1] [,2] [,3]
38 [1,]    1    4    7
39 [2,]    2    5    8
40 [3,]    3    6    9
41
42 $M2
43      [,1]
44 [1,]    9
45 [2,]    8
46 [3,]    7
47 [4,]    6
48 [5,]    5
49 [6,]    4
50 [7,]    3
51 [8,]    2
52 [9,]    1
53
54 $M3
55      [,1] [,2] [,3]
56 [1,]   11   14   17
57 [2,]   12   15   18
58 [3,]   13   16   19
59 > sapply(X = L3, FUN = apply, 2, median)
60 $M1
61 [1] 2 5 8
62
63 $M2
64 [1] 5
65
66 $M3
67 [1] 12 15 18

```

Oproti funkci `lapply()` jsou parametry `simplify` (výchozí `true`) a `USE.NAMES` (takéž) nové. Parametr `simplify` právě způsobí převod na pole namísto listu. Parametr `USE.NAMES` slouží v případě, kdy `X` je charakter, a způsobí použití `X` jako jména objektu.

4.2.3 Funkce vapply

Funkce `vapply()` je verzí funkce `sapply()` s tím rozdílem, že oproti `sapply()` je nutné definovat typ objektu, který se má navrátit. To, že definujeme typ vektoru (atomický typ elementů vektoru/pole, co se navrácí), má nespornou výhodu při kontrole toho, zda funkce proběhla správně – tedy, že navrácí hodnoty, jaké má. Další výhodou funkce oproti `sapply()` je navíc to, že je vykonána znatelně rychleji než funkce `sapply()`. Z těchto důvodů bývá funkce `vapply()` doporučována jako funkce, která se má používat spíše než `sapply()`. Použití funkce `vapply()` je pak následující:

```
1 > L <- list(
2 +   a = c(1:10),
3 +   b = c(2:5),
4 +   c = c(-5:5),
5 +   M = matrix(data = c(1:9), ncol = 3)
6 + )
7 > print(L)
8 $a
9  [1]  1  2  3  4  5  6  7  8  9 10
10
11 $b
12 [1]  2  3  4  5
13
14 $c
15 [1] -5 -4 -3 -2 -1  0  1  2  3  4  5
16
17 $M
18      [,1] [,2] [,3]
19 [1,]    1    4    7
20 [2,]    2    5    8
21 [3,]    3    6    9
22
23 > vapply(X = L,
24 +       FUN = mean,
25 +       FUN.VALUE = numeric(1))
26   a    b    c    M
27 5.5 3.5 0.0 5.0
28
29 > vapply(X = L,
30 +       FUN = mean,
31 +       FUN.VALUE = complex(1))
32   a      b      c      M
33 5.5+0i 3.5+0i 0.0+0i 5.0+0i
34
35 > vapply(X = L,
36 +       FUN = mean,
```



```

37 + FUN.VALUE = character(1))
38 Error in vapply(X = L, FUN = mean, FUN.VALUE = character(1)) :
39   values must be type 'character',
40   but FUN(X[[1]]) result is type 'double'

```

Parametr `FUN.VALUE` pak specifikuje typ proměnné, kterou má funkce `vapply()` navrátit. Zatímco `double` je interpretovatelná jako numerická i jako komplexní číslo, tak ji naopak jako charakter nelze navrátit. Již víme, že lze číslo převést na charakter pomocí funkce `as.character()`, ale zde jde o to, že návratová hodnota funkce `mean` není v souladu s očekávanou hodnotou zadanou uživatelem.

4.2.4 Funkce `replicate`

Funkce `replicate()` replikuje aplikaci funkce jako takové. Tedy neprovádí se zde pro prvky určitého objektu jedna funkce, ale primárně se provádí replikace jedné a té samé funkce. Ukažme si použití této funkce na funkci `sample()` v následujícím příkladu.

```

1 > set.seed(1)
2 > replicate(n = 5, exp = sample(x = 1:100, size = 3))
3      [,1] [,2] [,3] [,4] [,5]
4 [1,]   27   91   95    7   69
5 [2,]   37   20   66   21   39
6 [3,]   57   89   62   18   76

```

Funkce navrácí výstup ve formátu matice, pokud je to možné. Parametr `n` určuje počet replikací dané funkce. Funkce `replicate()` je podobná funkci cyklu `for` s tím, že zde není přítomný iterátor. Funkce má zejména použití v simulacích a v uživateli definovanými funkcemi.

Příklad 3

Máte následující list `L`.

```

1 > L <- list(
2 +   L1 = sample(1:100, 10),
3 +   L2 = sample(1:100, 5),
4 +   L3 = sample(1:100, 9)
5 + )
6 > print(L)
7 $L1
8 [1] 19 82 66 78 11 69 39 77 60 72
9
10 $L2
11 [1] 56 53 78 3 46

```

```
12
13 $L3
14 [1] 74 69 47 84 43 24 7 10 30
```

Pro tento list chcete zjistit délky vektorů, které se v něm nacházejí. Zkuste to provést funkcemi `sapply()`, `lapply()` a `vapply()`.

Řešení 3

```
1 > lapply(L, FUN = length)
2 $L1
3 [1] 10
4
5 $L2
6 [1] 5
7
8 $L3
9 [1] 9
10
11 > sapply(L, FUN = length)
12 L1 L2 L3
13 10 5 9
14 > vapply(L, FUN = length, FUN.VALUE = numeric(1))
15 L1 L2 L3
16 10 5 9
```

4.3 Funkce mapply

Zvláštním případem `apply()` funkcí je funkce `mapply()`. Jedná se o verzi funkce `sapply()` spojenou s namapováváním parametrů. Funkce je zejména vhodná, když potřebujeme vyzkoušet různé varianty verze nějakého zpracování aj. Jako standardní příklad pro funkci `mapply()` se uvádí použití na funkci `rep()`. Tato funkce má dva hlavní parametry: `x` (to, co se má replikovat) a `times` (kolikrát se to má replikovat).

```
1 > mapply(rep, x = 1:3, times = c(2, 4, 6))
2 [[1]]
3 [1] 1 1
4
5 [[2]]
6 [1] 2 2 2 2
7
8 [[3]]
9 [1] 3 3 3 3 3 3
```

Funkce `mapply()` nejen „namapuje“ prvky dvou vektorů na sebe, ale také i objekty v listech. Namapování zde pak probíhá prostřednictvím jednotlivých elementů listů. V následujícím případě na sebe napojíme ale list a vektor. První element listu (neatomický vektor) s prvním elementem vektoru (prvek o jedné hodnotě).

```
1 > mapply(rep, x = list(1:3, 3:4, 4:5), times = c(2, 3, 3))
2      [,1] [,2] [,3]
3 [1,]    1    3    4
4 [2,]    2    4    5
5 [3,]    3    3    4
6 [4,]    1    4    5
7 [5,]    2    3    4
8 [6,]    3    4    5
```

```
1 > data1 <- data.frame(A1 = c(1, 2, 1, 3),
2 +                   B1 = c(1, 2, 3, 4),
3 +                   C1 = c(1, 1, 2, 1))
4 > data2 <- data.frame(A2 = c(1, 2, 3, 4),
5 +                   B2 = c(1, 2, 2, 5),
6 +                   C2 = c(2, 2, 3, 1))
7 > print(data1)
8   A1 B1 C1
9  1  1  1  1
10 2  2  2  1
11 3  1  3  2
12 4  3  4  1
13 > print(data2)
14  A2 B2 C2
15  1  1  1  2
16  2  2  2  2
17  3  3  2  3
18  4  4  5  1
19 > mapply(cor, data1, data2)
20      A1      B1      C1
21 0.6741999 0.8944272 0.8164966
22 > # analogicky vystup by slo ziskat trojim pouziti funkce cor:
23 > cor(data1$A1, data2$A2)
24 [1] 0.6741999
25 > cor(data1$B1, data2$B2)
26 [1] 0.8944272
27 > cor(data1$C1, data2$C2)
28 [1] 0.8164966
```

Příklad 4

Jak byste pomocí funkce `mapply()` vytvořili následující matici?

```
1      [,1] [,2] [,3] [,4] [,5]
2 [1,]    1    2    3    4    5
3 [2,]    2    3    4    5    6
4 [3,]    3    4    5    6    7
```

Řešení 4

Stačí funkci `mapply()` použít spolu s funkcí `seq()`.

```
1 > mapply(seq, from = 1:5, to = 3:7)
2      [,1] [,2] [,3] [,4] [,5]
3 [1,]    1    2    3    4    5
4 [2,]    2    3    4    5    6
5 [3,]    3    4    5    6    7
```

4.4 Funkce `rapply`

Funkce `rapply()` je rekurzivní verzí `lapply()` funkce. Výhodou funkce je i parametr `how`, který nám umožňuje definovat, jak má funkce iterovat. Tento parametr má tři základní možnosti, jak se mají data párovat: `unlist`, `replace` a `list`. Dále umožňuje definovat typ dat, pro která se má funkce provést (atribut `class`).

```
1 > L <- list(
2 +   a = c(1:10),
3 +   b = c(2:5),
4 +   c = c(-5:5),
5 +   M = matrix(data = c(1:9), ncol = 3)
6 + )
7 > print(L)
8 $a
9  [1]  1  2  3  4  5  6  7  8  9 10
10
11 $b
12 [1] 2 3 4 5
13
14 $c
15 [1] -5 -4 -3 -2 -1  0  1  2  3  4  5
16
17 $M
18      [,1] [,2] [,3]
```

```

19 [1,]    1    4    7
20 [2,]    2    5    8
21 [3,]    3    6    9
22 > rapply(L, sum, how = "unlist")
23  a  b  c  M
24 55 14  0 45
25 > rapply(L, sum, how = "replace")
26 $a
27 [1] 55
28
29 $b
30 [1] 14
31
32 $c
33 [1] 0
34
35 $M
36 [1] 45
37 > rapply(L, sum, how = "list")
38 $a
39 [1] 55
40
41 $b
42 [1] 14
43
44 $c
45 [1] 0
46
47 $M
48 [1] 45

```

Hlavní rozdíl mezi funkcí `apply` a funkcí `rapply` spočívá v tom, že u funkce `rapply` je nejprve objekt vyhodnocen (na splnění podmínek – např. typu objektu, pro který se má zpracovat) a až následně je volaná samotná funkce v jazyce C. Nebuďte zde překvapeni, že se jedná o jazyk C. Software **R** je naprogramován v jiném jazyce než jazyk **R**. Většina základních funkcí je z důvodu optimalizace naprogramována v jazyce C, protože je výrazně rychlejší než samotný jazyk **R**. Do verze jazyka **R** 3.5.X bylo požadováno, aby vstupním objektem byl typ `list`, což již u novějších instalací není povinné a může jít i o tzv. list-like objekt či objekt typu „expression“, díky čemuž jsme schopni tuto funkci použít i pro `data.frame` objekt. Ukažme si toto na příkladu datasetu `iris`. Tento dataset obsahuje údaje o měření provedeném na různých druzích kosatců. Dataset obsahuje čtyři numerické proměnné a jednu kategoriální proměnnou uvažovanou jako typ objektu faktor (druh kosatce).

```

1 > head(iris)
2   Sepal.Length Sepal.Width Petal.Length Petal.Width Species
3 1           5.1         3.5         1.4         0.2   setosa

```

```

4 2          4.9          3.0          1.4          0.2  setosa
5 3          4.7          3.2          1.3          0.2  setosa
6 4          4.6          3.1          1.5          0.2  setosa
7 5          5.0          3.6          1.4          0.2  setosa
8 6          5.4          3.9          1.7          0.4  setosa
9 > rapply(iris, mean, classes = "numeric")
10 Sepal.Length Sepal.Width Petal.Length Petal.Width
11      5.843333      3.057333      3.758000      1.199333
12
13 >rapply(iris, summary, classes = "factor")
14      Species.setosa Species.versicolor Species.virginica
15              50              50              50

```

4.5 Funkce tapply a by funkce

Funkce `tapply()` je specifickou funkcí z rodiny `apply`, protože se funkce `tapply()` používá na vektor `dat`. K tomuto vektoru `dat` ale dodáváme vektor tříd (obvykle faktorů), podle kterých chceme data ze vstupního vektoru rozdělit pro výpočet. Představme si tedy, že máme vektor pozorování příslušících k různým skupinám faktorů (například měření podle druhů rostlin) a chceme spočítat statistiky za dané skupiny zvlášť (druhy rostlin). Ukažme si dva příklady použití této funkce na datech `iris` a `mtcars`. V prvním případě spočítáme pro proměnnou `Sepal.Length` průměr pro jednotlivé druhy rostlin.

```

1 > tapply(X = iris$Sepal.Length,
2 +       INDEX = iris$Species,
3 +       FUN = mean)
4      setosa versicolor  virginica
5      5.006      5.936      6.588

```

Funkce `tapply()` nám srozumitelně navrací pojmenované pole (jedorozměrné pole) dat příslušných průměrů pro jednotlivé druhy rostlin. Co kdybychom si chtěli vypočítat průměrnou vzdálenost, kterou ujedou auta na jeden galon benzínu (`mpg`) podle druhu převodovky (`am`) z datasetu `mtcars`?

```

1 > tapply(X = mtcars$mpg,
2 +       INDEX = mtcars$am,
3 +       FUN = mean)
4      0      1
5 17.14737 24.39231

```

Výhodou parametru `INDEX` je i to, že umí zpracovat více proměnných, podle kterých chceme výstup třídit. Můžeme si tedy spočítat průměrnou vzdálenost, kterou ujedou daná auta na galon benzínu podle druhu převodovky a počtu válců auta.

```

1 > > tapply(
2 +   X = mtcars$mpg,
3 +   INDEX = list(mtcars$am, mtcars$cyl),
4 +   FUN = mean
5 + )
6           4           6           8
7 0 22.900 19.12500 15.05
8 1 28.075 20.56667 15.40

```

Z výstupu je pak hezky vidět, že druh převodovky je důležitý z pohledu spotřeby auta, ale jeho rozdíl mezi skupinami klesá při nárůstu počtu válců.

Funkce `tapply()` se provádí pro jeden vektor. Co kdybychom ale chtěli spočítat data za všechny proměnné? Máme dvě možnosti. Buď použijeme funkci `tapply()` uvnitř funkce `apply()`, nebo použijeme tzv. wrapper funkci – `by()`. Zkusme nejprve „vložit“ funkci `tapply()` do funkce `apply()`. První problém, na který narazíme, je to, že funkce mají stejně pojmenované parametry. Toto lze vyřešit například tím, že parametr `FUN` pro funkci `tapply()` předáme nepojmenovaný.

```

1 > apply(
2 +   X = iris[, 1:4],
3 +   MARGIN = 2,
4 +   FUN = tapply,
5 +   INDEX = iris$Species,
6 +   mean
7 + )
8           Sepal.Length Sepal.Width Petal.Length Petal.Width
9 setosa           5.006         3.428         1.462         0.246
10 versicolor       5.936         2.770         4.260         1.326
11 virginica        6.588         2.974         5.552         2.026

```

Identický výsledek, avšak ve formě jiného objektu (objekt typu `by`), lze získat použitím funkce `by()`. Zkuste se teď podívat na obě řešení a najít zásadní rozdíly.

```

1 > by(data = iris[, 1:4],
2 +   INDICES = iris$Species,
3 +   FUN = colMeans)
4 iris$Species: setosa
5 Sepal.Length Sepal.Width Petal.Length Petal.Width
6           5.006         3.428         1.462         0.246
7 -----
8 iris$Species: versicolor
9 Sepal.Length Sepal.Width Petal.Length Petal.Width
10           5.936         2.770         4.260         1.326
11 -----
12 iris$Species: virginica

```

13	Sepal.Length	Sepal.Width	Petal.Length	Petal.Width
14	6.588	2.974	5.552	2.026

Je zřejmé, že prvním rozdílem ve způsobu provedení jsou rozdílné typy výstupních objektů. Funkce `apply()` nám navrátila matici, se kterou se dá dobře pracovat. S objektem typu `by` lze také pracovat dál, ale není to přímočaré (například převést na list, dále provést `unlist()` a nebo převést na vektor přímo). Za hlavní rozdíl lze ale považovat to, že funkce `tapply()` používá funkci `mean()` a funkce `by()` funkci `colMeans()`. S funkcí `mean()` by funkce `by()` vůbec nefungovala. To je dáno tím, že funkce `by()` zpracovává vstupní objekt jako celek. Pokud vám funkce `by` nevyhovuje a přijde vám uživatelsky příjemnější použití funkce `tapply()` v `apply()` funkci, tak mám pro vás dobrou zprávu: funkce `aggregate()` je řešením této situace.

Příklad 5

Jak byste pomocí funkce `tapply()` spočítali průměrnou hodnotu proměnné `mpg` z datasetu `mtcars` v třídění pro auta podle toho, zda mají podprůměrnou anebo nadprůměrnou proměnnou `hp`?

Řešení 5

```
1 > tapply(mtcars$mpg,
2 +       INDEX = mtcars$hp > mean(mtcars$hp),
3 +       FUN = mean)
4   FALSE      TRUE
5 24.22353 15.40667
```

Příklad 6

Chcete zjistit počet aut podle počtu válců a podle toho, zda mají podprůměrnou a nadprůměrnou hodnotu `hp`.

Řešení 6

```
1 > tapply(mtcars$mpg,
2 +       INDEX = list(mtcars$hp > mean(mtcars$hp), mtcars$cyl),
3 +       FUN = length)
4       4 6 8
5 FALSE 11 6 NA
6 TRUE  NA 1 14
```


Příklad 7

Dokázali byste úlohu z příkladu 5 provést pro všechny proměnné v datasetu?

Řešení 7

```
1 > by(mtcars,
2 +   INDICES = as.factor(mtcars$hp > mean(mtcars$hp)),
3 +   FUN = colMeans)
4 as.factor(mtcars$hp > mean(mtcars$hp)): FALSE
5      mpg      cyl      disp      hp
6 24.2235294  4.7058824 134.9823529  93.5294118
7      drat      wt      qsec      vs
8  3.8976471  2.5995882  18.8735294   0.8235294
9      am      gear      carb
10 0.5882353  3.9411765  2.0588235
11 -----
12 as.factor(mtcars$hp > mean(mtcars$hp)): TRUE
13      mpg      cyl      disp      hp
14 15.4066667  7.8666667 339.2266667 206.9333333
15      drat      wt      qsec      vs
16  3.2553333  3.917267  16.687333  0.0000000
17      am      gear      carb
18 0.2000000  3.400000  3.6666667
```

Příklad 8

Zvládli byste vypočítat korelační matici pro numerické proměnné v datasetu `iris` tříděného podle druhu rostliny (`Species`) pomocí funkce `by()`?

Řešení 8

```
1 > by(iris[, 1:4], INDICES = iris$Species, FUN = cor)
2 iris$Species: setosa
3      Sepal.Length Sepal.Width Petal.Length Petal.Width
4 Sepal.Length      1.0000000    0.7425467    0.2671758    0.2780984
5 Sepal.Width      0.7425467    1.0000000    0.1777000    0.2327520
6 Petal.Length      0.2671758    0.1777000    1.0000000    0.3316300
7 Petal.Width      0.2780984    0.2327520    0.3316300    1.0000000
8 -----
9 iris$Species: versicolor
10      Sepal.Length Sepal.Width Petal.Length Petal.Width
11 Sepal.Length      1.0000000    0.5259107    0.7540490    0.5464611
```

```

12 Sepal.Width      0.5259107      1.0000000      0.5605221      0.6639987
13 Petal.Length     0.7540490      0.5605221      1.0000000      0.7866681
14 Petal.Width      0.5464611      0.6639987      0.7866681      1.0000000
15 -----
16 iris$Species: virginica
17      Sepal.Length Sepal.Width Petal.Length Petal.Width
18 Sepal.Length     1.0000000     0.4572278     0.8642247     0.2811077
19 Sepal.Width      0.4572278     1.0000000     0.4010446     0.5377280
20 Petal.Length     0.8642247     0.4010446     1.0000000     0.3221082
21 Petal.Width      0.2811077     0.5377280     0.3221082     1.0000000

```

4.6 Funkce aggregate a ts objekt

Funkce `aggregate()` je agregační funkce, která umožňuje pro různé typy objektů vykonávat různé úlohy. Funkce umí, tak jako funkce `lapply()`, vykonat určitou funkci na vektoru, který třídíme do určitých skupin. Výhodou funkce `aggregate()` je, že výstupem v tomto případě je objekt typu `data.frame`. Jako hlavní parametr vložíme požadovaný vektor a do parametru `by` vložíme list, podle kterého chceme vykonat funkci (parametr `FUN`).

```

1 > aggregate(mtcars$mpg, by = list(mtcars$am), FUN = mean)
2   Group.1      x
3 1      0 17.14737
4 2      1 24.39231

```

Výstupem funkce `aggregate()` je vždy `data.frame`. V případě více klasifikací objektu (kategorií, do kterých prvky objektu třídíme), pro který danou úlohu vykonáváme, nám funkce `lapply()` navrací vícerozměrné pole, kde každý rozměr odpovídá dané klasifikaci dat (`am`, `gear`...). V případě funkce `aggregate()` nám funkce ale vždy navrátí `data.frame`. V případě vícerozměrné úlohy je výstupní objekt stále `data.frame`, ale přibude zde další sloupec proměnné, podle kterého identifikujeme výslednou statistiku ve sloupci `x`.

```

1 > aggregate(mtcars$mpg,
2           by = list(mtcars$am, mtcars$cyl),
3           FUN = mean)
4   Group.1 Group.2      x
5 1      0      4 22.90000
6 2      1      4 28.07500
7 3      0      6 19.12500
8 4      1      6 20.56667
9 5      0      8 15.05000
10 6      1      8 15.40000

```

Čtenáři doporučuji srovnat tento výstup s výstupem funkce `tapply()`. Hlavní potenciál funkce `aggregate()` ale spočívá v použití na celý objekt jako takový. Toto použití je pro typ objektu `data.frame`, ale výpočet se provede pro jakýkoliv objekt, který lze převést na `data.frame` pomocí funkce `as.data.frame()`.

```
1 > aggregate(mtcars ,
2 +           by = list(mtcars$gear),
3 +           FUN = mean)
4   Group.1      mpg      cyl      disp      hp      drat
5 1         3 16.10667  7.466667 326.3000 176.1333 3.132667
6 2         4 24.53333  4.666667 123.0167  89.5000 4.043333
7 3         5 21.38000  6.000000 202.4800 195.6000 3.916000
8           wt      qsec      vs      am gear      carb
9 1 3.892600 17.692 0.2000000 0.0000000      3 2.666667
10 2 2.616667 18.965 0.8333333 0.6666667      4 2.333333
11 3 2.632600 15.640 0.2000000 1.0000000      5 4.400000
```

Stejně jako do jakékoliv funkce z rodiny `apply()`, tak i do funkce `aggregate()` můžeme přiřadit anonymní funkci. O tomto se více dozvíte v kapitole Uživatelské funkce. Zde ale uvedeme jednoduchý příklad.

```

1 > aggregate(
2 +   x = iris[, 1:4],
3 +   by = list(iris$Species),
4 +   FUN = function(x)
5 +     sd(x) / mean(x)
6 + )
7      Group.1 Sepal.Length Sepal.Width Petal.Length Petal.Width
8 1      setosa   0.07041344   0.1105789   0.11878522   0.4283967
9 2 versicolor   0.08695606   0.1132846   0.11030774   0.1491348
10 3  virginica   0.09652089   0.1084387   0.09940466   0.1355627

```

Pro datový soubor `iris` jsme spočítali variační koeficient. Tedy směrodatnou odchylku dělenou průměrem. **R** nemá funkci, která by toto přímo spočítala v základní instalaci. Proto jsme si do parametru `FUN` dosadili postup výpočtu, který chceme pro dané sloupce provést. Více se o této problematice dozvíte ale až v kapitole věnované tvorbě funkcí. Jediná nevýhoda funkce `aggregate()` je pak v tom, že oproti funkci `by` nezvládne zpracovat maticový výstup (viz příklad 8).

Speciální použití funkce `aggregate()` pak souvisí s objektem typu časové řady. Časové řady vytváříme v jazyku **R** pomocí funkce `ts`. Tato funkce má jako vstupní parametr `data` (vektor a nebo i `data.frame` objekt), dále parametr `start` (počáteční období) a `end` (poslední záznam). Další důležitý parametr je četnost pozorování na jednotku času – `frequency`. Jedná se o to, s jakou periodicitou je daná časová řada pozorována. V případě čtvrtletní časové řady volíme `frequency = 4`, v případě měsíční `frequency = 12`. Vytvoříme si měsíční časovou řadu z vektoru `y`, který nabývá hodnot jedna až sto dvacet. A řekněme, že první pozorování je za leden roku 2000.

```

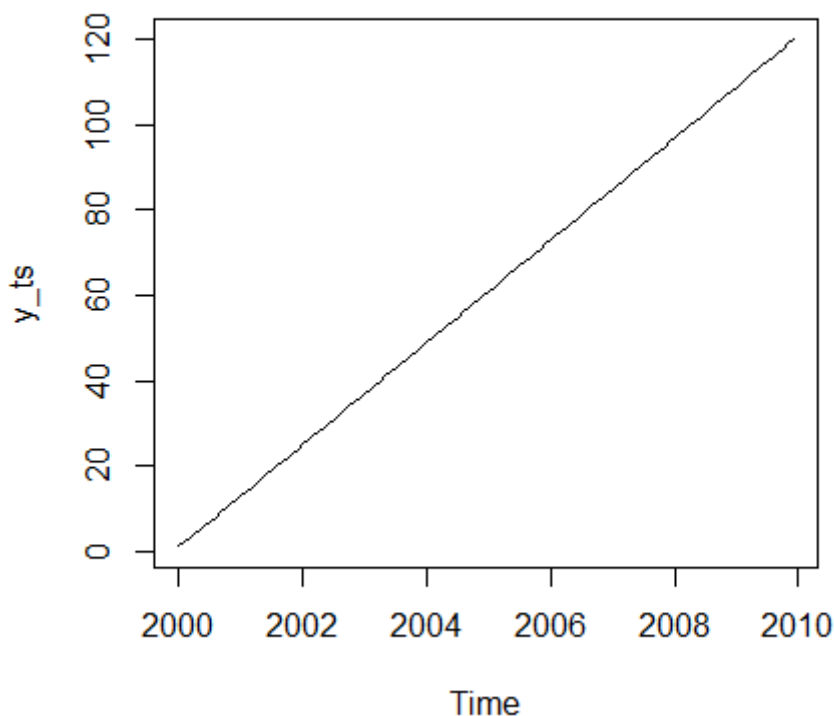
1 > y = c(1:120)
2 > y_ts <- ts(data = y,
3 +           start = c(2000, 1),
4 +           frequency = 12
5 > y_ts
6      Jan Feb Mar Apr May Jun Jul Aug Sep Oct Nov Dec
7 2000    1  2   3   4   5   6   7   8   9  10  11  12
8 2001   13  14  15  16  17  18  19  20  21  22  23  24
9 2002   25  26  27  28  29  30  31  32  33  34  35  36
10 2003   37  38  39  40  41  42  43  44  45  46  47  48
11 2004   49  50  51  52  53  54  55  56  57  58  59  60
12 2005   61  62  63  64  65  66  67  68  69  70  71  72
13 2006   73  74  75  76  77  78  79  80  81  82  83  84
14 2007   85  86  87  88  89  90  91  92  93  94  95  96
15 2008   97  98  99 100 101 102 103 104 105 106 107 108
16 2009  109 110 111 112 113 114 115 116 117 118 119 120

```

V případě, že využijeme parametru `start`, již nemusíme použít parametr `end` a naopak. **R**

má pro objekt `ts` předpřipravené chování několika základních funkcí. Podívejme se, co se stane, když na objekt `ts` zavoláme funkci `plot`.

```
1 plot(y_ts)
```



Obrázek 4.1: Plot funkce

Zdroj: Vlastní zpracování.

Funkce `plot()` má předprogramované procedury pro objekty typu `ts` a vykreslí je předem předpřipraveným způsobem. O tom, jak se dělají a upravují grafy, si však povíme více až v následujících kapitolách. U časových řad často řešíme situaci, kdy chceme agregovat data. Pro tento případ má funkce `aggregate()` parametr `nfrequency` neboli „new frequency“. Do tohoto parametru specifikujete četnost nové časové řady. Ukážeme si to na následujícím příkladu.

```
1 > aggregate(x = y_ts, nfrequency = 4)
2      Qtr1 Qtr2 Qtr3 Qtr4
```

3	2000	6	15	24	33
4	2001	42	51	60	69
5	2002	78	87	96	105
6	2003	114	123	132	141
7	2004	150	159	168	177
8	2005	186	195	204	213
9	2006	222	231	240	249
10	2007	258	267	276	285
11	2008	294	303	312	321
12	2009	330	339	348	357

Kam ale zmizel parametr FUN? V případě, že vstupním objektem je časová řada, má parametr FUN přiřazenou výchozí hodnotu na funkci sum. Automaticky se tedy provede součet. V případě, že chceme provést pro daná čtvrtletí jiný výpočet, je nutné specifikovat parametr FUN. Chceme-li například spočítat průměr za daná čtvrtletí, tak toho dosáhneme pomocí následujícího příkladu.

```

1 > aggregate(x = y_ts,
2 +           nfrequency = 4,
3 +           FUN = mean)
4           Qtr1 Qtr2 Qtr3 Qtr4
5 2000         2    5    8    11
6 2001        14   17   20   23
7 2002        26   29   32   35
8 2003        38   41   44   47
9 2004        50   53   56   59
10 2005        62   65   68   71
11 2006        74   77   80   83
12 2007        86   89   92   95
13 2008        98  101  104  107
14 2009       110  113  116  119

```

Zkusme si toto použití ukázat na reálných datech. V R se kromě datasetů mtcars a iris nachází také dataset AirPassengers. Jedná se o měsíční počty pasažérů mezi roky 1949 až 1960 v tisících v letecké dopravě. Zkusme si pro ně vypočítat variabilitu v jednotlivých čtvrtletích pomocí funkce sd.

```

1 > aggregate(x = AirPassengers,
2 +           nfrequency = 4,
3 +           FUN = sd)
4           Qtr1      Qtr2      Qtr3      Qtr4
5 1949 10.263203  7.023769  6.928203  8.386497
6 1950 13.051181 12.055428  6.928203 13.453624
7 1951 17.785762  7.549834  8.660254 10.583005
8 1952 11.060440 20.808652 16.703293 11.930353
9 1953 23.094011  7.023769 18.339393 15.821926

```

```

10 1954 23.895606 19.655364 22.678918 15.011107
11 1955 17.616280 26.274195 26.514147 22.605309
12 1956 21.361960 33.867388 31.432467 20.207259
13 1957 28.583212 40.853396 35.809682 21.779195
14 1958 22.000000 46.508064 54.720502 24.542480
15 1959 33.005050 38.850139 52.538874 25.423087
16 1960 15.620499 39.929104 61.719797 35.697806

```

Příklad 9

Zkuste, jako v příkladu č. 7, vypočítat průměrnou hodnotu jednotlivých proměnných v datasetu `mtcars` tříděného podle toho, zda auto má anebo nemá nadprůměrnou hodnotu proměnné `hp`.

Řešení 9

```

1 > aggregate(mtcars,
2 + by = list(mtcars$hp > mean(mtcars$hp)), FUN = mean)
3   Group.1      mpg      cyl      disp      hp      drat      wt
4 1   FALSE 24.22353  4.705882 134.9824  93.52941  3.897647 2.599588
5 2    TRUE 15.40667  7.866667 339.2267 206.93333  3.255333 3.917267
6
7   qsec      vs      am      gear      carb
8 1 18.87353 0.8235294 0.5882353  3.941176  2.058824
9 2 16.68733 0.0000000 0.2000000  3.400000  3.666667

```

Kapitola 5

Tvorba vlastních funkcí

5.1 Uživatelem definované funkce

Funkci si můžeme představit jako určitou část kódu, kterou někdo zabalil a upravil tak, abychom mohli zadat určité parametry (argumenty) a funkce nám navrátila výsledek. Funkce obvykle interaguje s programovacím prostředím pomocí parametrů funkce a návratové hodnoty funkce. Návratovou hodnotou se rozumí hodnota (neboli výsledek), kterou funkce navrácí (více viz. Wickham, 2020a; Long, Teetor, 2019; R studio team, 2018).

Funkce je z teoretického pohledu podprogramem, který umožňuje abstraktnější práci s kódem. Výhodou tvorby funkcí je, že kód lze zpřehlednit a mnoho operací i zautomatizovat. Uživatel jazyka **R** by měl brát na zřetel, že funkce jsou objekty jako jiné v jazyku **R**. S uživatelem definovanými funkcemi lze pracovat jako s jinými objekty – například se dají odstranit. Uživatelem definovaná funkce se uloží po boku ostatních definovaných objektů a lze si je zobrazit pomocí funkce `ls()`. Následující kód ilustruje to, jak taková definice funkce vypadá a z čeho se skládá.

```
1 Jmeno_funkce <- function(#parametry_funkce){  
2   # Telo_funkce  
3 }
```

Uživatelem definovaná funkce („user defined function“) má následující části:

- jméno,
- parametry funkce neboli parametry,
- tělo funkce (kód, který se realizuje),
- návratovou hodnotu (je součástí těla funkce).

Ukažme si v praxi velmi jednoduchou funkci, která pouze sečte dvě čísla.

```
1 > secist <- function(x, y) {  
2 +   x + y  
3 + }
```

Nejprve je nutné funkci vyhodnotit neboli zavést (tedy spustit v příkazové řádce). Příkazová řádka při spuštění kódu funkce nic nenavráť. Jedná se jen o inicializaci funkce a není třeba (v případě, že nenastala chyba) cokoliv do příkazové řádky navracet. To, že funkce je k dispozici, si například ověříme tak, že zavoláme příkaz `ls` a podíváme se na objekty, které jsou dostupné.

```
1 > ls()  
2 [1] "secist"
```

Jak vidíte, v mém případě je funkce již vytvořená. Zobrazuje se v seznamu objektů. Funkci používáme tak, že napíšeme její jméno a do kulatých závorek pak parametry funkce.

```
1 > secist(x = 2, y = 3)  
2 [1] 5
```

V případě, kdy zapomeneme za jméno funkce napsat závorky a zavoláme jen jméno dané funkce, tak jako výsledek se nám do příkazové řádky navrátí celá samotná funkce – její kód.

```
1 > secist  
2 function(x, y) {  
3   x + y  
4 }
```

V následujících kapitolách se podíváme na jednotlivé aspekty tvorby funkcí v jazyku **R**.

5.1.1 Jméno funkce

U jmen funkcí, tak jako u jmen objektů, dodržujeme pravidlo, že jméno nesmí začínat číslicí a nesmí obsahovat speciální znak. Dále je nutné se vyhnout rezervovaným jménům. Ne každá uživatelem definovaná funkce ale musí mít jméno. Takovým případem je anonymní funkce. Tento typ funkce má speciální použití zejména v matematických aplikacích nebo ve funkcích typu `apply()`. Obvykle se používá v situaci, kdy se očekává, že funkce bude použita jen na jednom místě v kódu. Následující funkce se jmenuje `rozpeti()` a přebírá jediný parametr `x`, ze kterého spočítá rozpětí (maximum – minimum dat). Jedná se o uživatelem definovanou funkci, která má jméno.

```
1 > rozpeti <- function(x) {  
2 +   max(x) - min(x)  
3 + }
```

Alternativně si můžeme ukázat analogickou funkci přímo použitou ve funkci `apply`.

```
1 > apply(X = mtcars,
2 +       MARGIN = 2,
3 +       FUN = function(x) max(x) - min(x)
4 + )
5      mpg      cyl      disp      hp      drat      wt      qsec
6 23.500    4.000  400.900  283.000    2.170    3.911    8.400
7      vs      am      gear      carb
8 1.000    1.000    2.000    7.000
```

Samozřejmě lze funkci `rozpeti()` použít i ve funkci `apply()`. Musí být ale funkce `rozpeti` zavedena. Zavedením funkce se zde rozumí to, že daná funkce již existuje v našem prostředí. Funkce se stávají součástí uživatelského prostředí tak jako jiné objekty a můžeme si je zobrazit pomocí `ls` funkce.

```
1 > apply(X = mtcars, MARGIN = 2, FUN = rozpeti)
2      mpg      cyl      disp      hp      drat      wt      qsec
3 23.500    4.000  400.900  283.000    2.170    3.911    8.400
4      vs      am      gear      carb
5 1.000    1.000    2.000    7.000
```

5.1.2 Parametry funkce

U každé funkce rozlišujeme tři druhy parametrů funkce. Jedná se o parametry funkce bez výchozí hodnoty, s výchozí hodnotou a dodatečné parametry. Funkce bez výchozích parametrů jsou ty, do kterých musí být přiřazena hodnota (pokud nastane situace, že se mají v tělu funkce použít). Parametry s výchozí hodnotou jsou ty, do kterých je přiřazena výchozí hodnota, a tedy uživatel nemusí vkládat nějakou hodnotu. Uživatel má ale možnost výchozí hodnotu nahradit svou vlastní hodnotou. Toto se provede prostým přiřazením nové hodnoty do daného parametru při aplikaci funkce – tak jako v případě povinných parametrů. Zajímavostí je, že výchozí hodnota nemusí být vyloženě hodnotou, ale může se jednat i o výpočet. Dokonce existují situace, kdy lze většinu výpočtu realizovat na úrovni parametrů funkce. Třetím typem parametrů jsou dodatečné parametry, tedy zda funkce může přijmout nějaké předem neznámé parametry. Podívejme se nejprve na příklad funkce s povinným parametrem, kdy pro `x` není vložena výchozí hodnota. Tato funkce vypočte variační rozpětí.

```
1 > rozpeti <- function(x) {
2 +   max(x) - min(x)
3 + }
4 > rozpeti(1:10)
5 [1] 9
```

Následující kód ilustruje použití funkce pro mocnění s přiřazenými výchozími hodnotami. Parametr `s` s výchozí hodnotou zadefinujeme pomocí jména daného parametru, rovnítka a hodnoty, které nabývá. Hodnoty přiřazené v rámci definice funkce nazýváme výchozí hodnoty, a pokud je uživatel nevloží do daných parametrů funkce, přebere dané hodnoty, jako by je definoval sám uživatel. Díky tomu, že funkce má pro oba parametry definovanou výchozí hodnotu, lze ji použít i bez jakýchkoli parametrů.

```
1 > sqxy <- function(x = 1, y = 0) {  
2 +   x ** y  
3 + }  
4 >  
5 > sqxy()  
6 [1] 1
```

V případě, kdy uživatel nezadá žádné hodnoty, bude hodnota `x` rovna jedné a `y` rovna nule. Alternativně lze pro oba parametry přiřadit hodnoty uživatelem. Uživatelem vložené hodnoty mohou být vložené v pořadí, v jakém je funkce uvádí, anebo pomocí parametrů funkce. Když využíváme přiřazení do parametrů dané funkce, tak nemusíme dodržet pořadí parametrů.

```
1 > sqxy(2, 3)  
2 [1] 8  
3  
4 > sqxy(x = 2, y = 3)  
5 [1] 8  
6  
7 > sqxy(2, y = 3)  
8 [1] 8  
9  
10 > sqxy(y = 3, x = 2)  
11 [1] 8
```

V případě, kdy chceme vynechat zadání parametru `x` a chceme jen definovat parametr `y`, tak musíme s ohledem na pořadí použít přiřazení hodnoty do příslušného parametru. A to proto, že parametr `x` je definovaný v deklaraci funkce jako první. Ukažme si tento jev v následujícím příkladu na situaci, kdy `y` je rovno třem.

```
1 > sqxy(y = 3)  
2 [1] 1
```

Všichni známe Pythagorovu rovnici pro pravoúhlý trojúhelník ($c^2 = a^2 + b^2$) ze středoškolského učiva. Následující kód ilustruje implementaci zmíněné rovnice pomocí funkce v jazyku **R**. Schválně jsme zde zvolili velmi specifickou implementaci. Následující příklad ilustruje funkci, jejíž celý výpočet se realizuje na úrovni definice parametrů funkce. Uživatel zadá hodnoty dvou parametrů a funkce zbylé hodnoty dopočítá – v samotné definici parametrů.

```

1 > pythagoras <- function(a = sqrt(c ^ 2 - b ^ 2),
2 +                       b = sqrt(c ^ 2 - a ^ 2),
3 +                       c = sqrt(a ^ 2 + b ^ 2)) {
4 +   list(a = a, b = b, c = c)
5 + }
6 > pythagoras(a = 1, c = 5)
7 $a
8 [1] 1
9
10 $b
11 [1] 4.898979
12
13 $c
14 [1] 5

```

Tento způsob definice parametrů funkce je velmi efektivní zejména v situaci, kdy předpokládáme, že uživatel je znalý dané problematiky (a tak nepředpokládáme, že by zadal sám všechny tři parametry, či že by zadal jen jeden). Daná definice má své výhody. Jedná se zejména o zefektivnění z pohledu rychlosti evaluace (vyhodnocení) dané funkce. Protože v případě, kdy bychom definovali výpočet hodnot chybějící strany trojúhelníku v těle funkce, pak by musela nad rámec výpočtu být spuštěna i `if` podmínka, která by danou situaci ověřila.

5.1.3 Tělo funkce

Tělo funkce představuje část funkce ve složených závorkách (popřípadě výraz následující za definicí parametrů funkce). V rámci těla funkce se realizuje samotný výpočet funkce. Jeho součástí může být i definice jiných funkcí.

```

1 > statistika <- function(x) {
2 +   prumer <- function(x) sum(x) / length(x)
3 +   sm.odch <- function(x) sqrt(var(x))
4 +   rozpeti <- function(x) max(x) - min(x)
5 +
6 +   list(
7 +     prumer = prumer(x),
8 +     sm.odchylka = sm.odch(x),
9 +     rozpeti = rozpeti(x)
10 +   )
11 + }
12
13 > statistika(1:10)
14 $prumer
15 [1] 5.5

```

```

16
17 $sm.odchylka
18 [1] 3.02765
19
20 $rozpeti
21 [1] 9

```

Výhodou této definice je to, že funkce nejsou dostupné ve standardním uživatelském prostředí. Jako všechny hodnoty definované uvnitř těla funkce, tak i tyto funkce po realizaci funkce mizí a nejsou nikde uloženy, tedy ani uživateli nejsou dostupné.

Prostředí funkce

Uživatelské prostředí obsahuje uživatelem definované proměnné a funkce. Můžeme si je zobrazit například pomocí funkce `ls`. Toto již známe z předcházejících kapitol. Uživatelské funkce se jmenují právě proto, že objekty v nich definuje uživatel. **R** má k dispozici ale více prostředí. Při realizaci jakékoliv funkce si daná funkce sama vytvoří vlastní prostředí s proměnnými. Toto prostředí není pro uživatele běžně přístupné a po tom, co funkce skončí, tak prostředí funkce zaniká. Ukažme si tuto situaci na následujícím příkladu.

```

1 > x <- 1
2 > print(x)
3 [1] 1
4 >
5 > moje_x <- function(y = 10) {
6 +   x <- y
7 +   list(x = x)
8 + }
9 >
10 > print(x)
11 [1] 1
12 > print(moje_x())
13 $x
14 [1] 10
15
16 > print(x)
17 [1] 1

```

Z předcházejícího příkladu je poměrně snadno vidět, že proměnná `x` zůstává proměnnou `x` i po tom, co je vyhodnocena funkce, která obsahuje sama definici stejně pojmenované proměnné. Důvodem, proč nedojde ke změně proměnné `x`, je to, že ta je funkcí vytvořena jen v rámci prostředí funkce, které navíc po vykonání funkce zaniká.

U parametrů funkce jsme schválně vynechali zmínku o tom, že funkce může přebírat i proměnné z nadřazeného prostředí. Tedy dokáže si „sáhnout“ pro proměnnou, která je nad

prostředím, ve kterém se funkce realizuje. Tuto situaci si můžeme ukázat na následujícím příkladu. Proměnná `y` není nijak použita v samotném těle funkce, a tak její výchozí hodnota nemá žádný dopad na samotný výsledek funkce. Co ale má dopad, je proměnná `x` definovaná v uživatelském prostředí.

```
1 > x <- 1
2
3 > moje_x <- function(y = 10) {
4 +   list(x = x)
5 + }
6 >
7 > moje_x()
8 $x
9 [1] 1
10
11 > x <- 10
12 > moje_x()
13 $x
14 [1] 10
```

To, že software a jazyk **R** postupuje ve vyhledávání dané proměnné postupně vždy do nadřazeného prostředí, ve kterém byla funkce definovaná, si můžeme hezky ukázat na následujících dvou příkladech. V tomto příkladu je `funkce_2()` vnořená do `funkce_1()`. Pokud `funkce_2()` nenajde v těle funkce určitou proměnnou, nejprve ji vyhledá v první funkci (`funkce_1()`) a pak až následně v uživatelském prostředí. Tedy přesto, že máme v uživatelském prostředí do hodnoty `x` přiřazenou hodnotu jedna, tak je navracena hodnota `funkce_1()`.

```
1 > x <- 1
2 > funkce_1 <- function() {
3 +   funkce_2 <- function()
4 +     list(x = x)
5 +   x <- "funkce_1"
6 +   funkce_2()
7 + }
8
9 > funkce_1()
10 $x
11 [1] "funkce_1"
```

Opakem k této situaci je pak situace, kdy funkce 2 je definována mimo funkci 1. Pak přesto, že ve `funkce_1()` je definována proměnná `x`, si `funkce_2()` „nesáhne“ na proměnnou `x` z `funkce_1()`, ale do uživatelského prostředí. Toto chování je způsobené tím, že `funkce_2()` je vyhodnocená v rámci `funkce_1()`, ale přitom sama je definována na stejné úrovni jako `funkce_1()`. Tedy v rámci uživatelského prostředí.

```

1 > x <- "uzivatelske prostredi"
2 > funkece_2 <- function() {
3 +   x
4 + }
5
6 > funkece_1 <- function() {
7 +   x <- "funkece_1"
8 +   list(x = funkece_2())
9 + }
10
11 > funkece_1()
12 $x
13 [1] "uzivatelske prostredi"
14
15 > funkece_2()
16 [1] "uzivatelske prostredi"

```

Lazy evaluation

Lazy evaluation je termín popisující postup vyhodnocení proměnných (a vlastně také to, zda proměnná vůbec existuje). Existence proměnné je ověřována, až když je volána samotným programem. Tedy dokud není potřebná při výpočtu, tak se o ni software nezajímá a je jedno, jestli obsahuje chybnou hodnotu, či zda vůbec existuje. Tuto situaci dobře ilustruje následující kód. Vyjdeme z předpokladu, že nemáte v prostředí jazyka definovanou proměnnou *Y*. Tato proměnná se ale používá v těle funkce. V případě, kdy se nevyhodnocuje část kódu, kde se daná proměnná používá, tak funkce proběhne bez chyby. Naopak, když dojde k realizaci části kódu s touto proměnnou, nastává chyba a výkon funkce se zastaví. Tímto postupem se v jazyku **R** vyhodnocují všechny funkce.

```

1 > mojefce <- function(x) {
2 +   if (x == 1) {
3 +     print(x)
4 +   } else{
5 +     print(Y)
6 +   }
7 + }
8
9 > mojefce(1)
10 [1] 1
11 > mojefce(3)
12 Error in print(Y) : object 'Y' not found

```

V předcházejícím případě není nikde zadefinovaná proměnná *Y* (ani na úrovni jiného prostředí), tak při druhém pokusu se navrací chybová hláška. Všimněte si ale, že v případě

prvního pokusu proběhne celá realizace kódu bez chyby. To je dáno tím, že se část kódu s proměnnou `Y` při splnění podmínky `x==1` nerealizovala.

5.1.4 Návrátová hodnota funkce

Každá funkce v jazyce **R** má návratovou hodnotu. Pokud není návratová hodnota definována, navrácí funkce hodnotu `NULL`. Z tohoto důvodu některé grafické funkce navrácí `NULL`, pokud se přiřadí do objektu. Návratová hodnota je standardně ta hodnota, která se vyhodnocovala (zavolala) jako poslední. V případě komplexních úloh může nastat situace, kdy si nejsme jisti, jaká hodnota byla poslední vyhodnocená. Z tohoto důvodu umožňuje jazyk **R** vynutit, jaký objekt funkce navrátí. Toto lze provést pomocí příkazu `return`, do něhož se jako parametr vloží objekt, který má daná funkce navrátit. Jak ale postupovat, pokud chceme navrátit několik možných objektů? To lze snadno vyřešit pomocí objektu typu `list`.

```
1 > mojefce <- function() {
2 +   return(list(a = 1, b = matrix(1:4, 2)))
3 +   x <- 10
4 +   x
5 + }
6 > mojefce()
7 $a
8 [1] 1
9
10 $b
11      [,1] [,2]
12 [1,]    1    3
13 [2,]    2    4
```

Příklad 1

Vytvořte funkci, která přebere jako parametr numerický vektor `dat` a provede jeho diferenci. Otestujte na vektoru `x<-1:10`.

Řešení 1

```
1 > difference <- function(x) {
2 +   return(x[2:(length(x))] - x[1:(length(x) - 1)])
3 + }
4
5 > x <- 1:10
6
7 > difference(x)
```



```
8 [1] 1 1 1 1 1 1 1 1 1 1
```

Příklad 2

Vytvořte funkci, která od vstupního vektoru dat odečte průměr a výsledek vydělí směrodatnou odchylkou. Tento matematický postup se nazývá standardizace dat, někdy také standardizované skóre (z-score). Použijte danou funkci na datech -2:2 a zkuste pro výsledná data vypočítat průměr a směrodatnou odchylku.

Řešení 2

```
1 > zscore <- function(x) {  
2 +   return((x - mean(x)) / sd(x))  
3 + }  
4 > # Výsledek funkce uložíme do proměnné y  
5 > # se kterou budeme pak dále pracovat  
6 > y <- zscore(-2:2)  
7 > print(y)  
8 [1] -1.2649111 -0.6324555  0.0000000  0.6324555  1.2649111  
9 > # průměr a směrodatná odchylka:  
10 > mean(y)  
11 [1] 0  
12 > sd(y)  
13 [1] 1
```

Příklad 3

Vytvořte tzv. wrapper funkci. Wrapper funkce je funkce, která obaluje nějakou funkci a automatizuje její použití na jiném datasetu. Obvykle jsme v situaci, kdy máme funkci, jejíž použití je určené na vektor, a chceme vytvořit wrapper funkci, která automatizuje její použití na objektu typu `data.frame`. Vytvořte wrapper funkci pro naši funkci `zscore`.

Řešení 3

```
1 > zscore <- function(x) {  
2 +   return((x - mean(x)) / sd(x))  
3 + }  
4  
5 > zscore_dataset <- function(x) {  
6 +   return(apply(x, 2, FUN = zscore))  
7 + }
```

```

8
9 > mtcars_st = zscore_dataset(mtcars)
10 > head(mtcars_st)
11
12      mpg      cyl      disp      hp
13 Mazda RX4      0.1508848 -0.1049878 -0.57061982 -0.5350928
14 Mazda RX4 Wag      0.1508848 -0.1049878 -0.57061982 -0.5350928
15 Datsun 710      0.4495434 -1.2248578 -0.99018209 -0.7830405
16 Hornet 4 Drive      0.2172534 -0.1049878  0.22009369 -0.5350928
17 Hornet Sportabout -0.2307345  1.0148821  1.04308123  0.4129422
18 Valiant      -0.3302874 -0.1049878 -0.04616698 -0.6080186
19
20      drat      wt      qsec      vs
21 Mazda RX4      0.5675137 -0.610399567 -0.7771651 -0.8680278
22 Mazda RX4 Wag      0.5675137 -0.349785269 -0.4637808 -0.8680278
23 Datsun 710      0.4739996 -0.917004624  0.4260068  1.1160357
24 Hornet 4 Drive -0.9661175 -0.002299538  0.8904872  1.1160357
25 Hornet Sportabout -0.8351978  0.227654255 -0.4637808 -0.8680278
26 Valiant      -1.5646078  0.248094592  1.3269868  1.1160357
27
28      am      gear      carb
29 Mazda RX4      1.1899014  0.4235542  0.7352031
30 Mazda RX4 Wag      1.1899014  0.4235542  0.7352031
31 Datsun 710      1.1899014  0.4235542 -1.1221521
32 Hornet 4 Drive -0.8141431 -0.9318192 -1.1221521
33 Hornet Sportabout -0.8141431 -0.9318192 -0.5030337
34 Valiant      -0.8141431 -0.9318192 -1.1221521

```

5.2 Anonymní funkce

Velmi důležitou roli hraje anonymní funkce. Anonymní funkce je funkce, která nemá přiřazené jméno. Tato funkce se používá zejména v situacích, kdy je nutné použít funkci, ale nechceme ji ukládat. Tato situace nastává často v kombinaci s jinými funkcemi, jako jsou funkce `apply` anebo funkce pro optimalizace či řešení soustav rovnic.

```

1 > x <- matrix(c(1, 3, 1, 5), 2)
2 > print(x)
3      [,1] [,2]
4 [1,]    1    1
5 [2,]    3    5
6 > apply(x, 2, FUN = function(x) max(x)-min(x))
7 [1] 2 4

```

Použití anonymní funkce má jeden vstup, a to parametr `x`. Za něj postupně funkce `apply()` dosazuje různé hodnoty – sloupce nebo řádky podle parametru `MARGIN` (druhý parametr funkce `apply`). Lze samozřejmě vytvořit i anonymní funkci s více parametry. V případě

použití ve funkci `apply()` je nutné dodatečně tyto parametry specifikovat za definicí anonymní funkce.

```
1 > x <- matrix(c(1, 3, 1, 5), 2)
2 > print(x)
3 [1] 1 2 3 4 5 6 7 8 9 10
4 > apply(x, 2, FUN=function(x, z) sum(x)*z, z=2)
5 [1] 8 12
```

Příklad 4

Pomocí anonymní funkce a funkce `apply` vypočítejte přes dataset `mtcars` variační koeficient.

Řešení 4

```
1 > apply(mtcars, MARGIN = 2, FUN = function(x) sd(x)/mean(x))
2      mpg      cyl      disp      hp      drat      wt
3 0.2999881 0.2886338 0.5371779 0.4674077 0.1486638 0.3041285
4      vs      am      gear      carb
5 1.1520369 1.2282853 0.2000825 0.5742933
```

5.3 Rekurzivní funkce

Specifickým typem funkce je tzv. rekurzivní funkce. Funkce je zvláštní tím, že volá sama sebe. Dobře si tuto funkci prohlédněte. Cílem této funkce je najít nejbližší celočíselně dělitelnou hodnotu. Tedy pokud chci najít nejbližší celočíselný dělitel čísla deset číslem tři, tak musím v první fázi zjistit, zda číslo deset je anebo není celočíselně dělitelné číslem tři beze zbytku. Pokud tomu tak není, hledám nižší číslo než číslo deset, které je celočíselně dělitelné třemi beze zbytku. V prvním kroku se tedy ověří, zda je číslo deset celočíselně dělitelné třemi beze zbytku. Jelikož je zbytek po dělení (`10 %% 3`) roven 1, tak funkce vstoupí do `else` části a sníží hodnotu `y` o jednu jednotku. Poté spustí znovu funkci nejbližšího dělitele s hodnotou `y` rovnou devíti. Hodnota devět má již zbytek po dělení roven nule a funkce navrátí hodnotu `y` jako výsledek (v našem případě devět).

```
1 > Nejblizsi_delitel <- function(y, x) {
2 +   if (y %% x == 0) {
3 +     return(y)
4 +   } else{
5 +     return(Nejblizsi_delitel(y - 1, x))
6 +   }
```

```

6 +   }
7 + }
8 >
9 > Nejblizsi_delitel(10, 3)
10 [1] 9
11 > Nejblizsi_delitel(115, 3)
12 [1] 114

```

V případě realizace `else` části podmínky dojde ke znovu zavolání té samé funkce. Tomuto se říká rekurze. Velkým problémem rekurzivních funkcí je ale situace, kdy počet rekurzí je příliš velký, popř. je nekonečný. V tu chvíli musíme do rekurze vložit podmínku, která kontroluje počet rekurzí. Aby počítadlo počtu rekurzí „prošlo“ do další funkce, je nutné přidat parametr funkce. Ve výchozí deklaraci je počet rekurzí nulový (`pocitadlo = 0`). Při realizaci `else` části kódu se ale k hodnotě počítadla rekurzí přičítá hodnota jedna.

```

1 > Nejblizsi_delitel <- function(y, x, pocitadlo = 0) {
2 +   if (pocitadlo < 100) {
3 +     if (y %% x == 0) {
4 +       return(y)
5 +     } else{
6 +       return(Nejblizsi_delitel(y - 1, x, pocitadlo + 1))
7 +     }
8 +   } else{
9 +     return("zastaveno z titulu velkeho mnozstvi rekurzi")
10 +   }
11 + }
12
13 > Nejblizsi_delitel(10, 3)
14 [1] 9
15 > Nejblizsi_delitel(10, 3.1)
16 [1] 0

```

Rekurze vede ke kódu, který je v některých případech velmi elegantní oproti kódu bez rekurze (je kratší a dobře čitelný). Nicméně když se ta samá situace vyřeší bez rekurze, je vždy trochu rychlejší (což samozřejmě nemusí být problém). Proto by se měla rekurze používat jen v případech, kdy je to opravdu vhodné.

Příklad 5

Přidejte do námi vytvořené funkce v příkladu č. 1 nový parametr jménem `d`. Tento parametr bude označovat řád difference. Vytvořte rekurzivní funkci, která spočítá diferenci daného řádu. Př. difference druhého řádu je pak difference z první difference. Aplikujte funkci při 1,2 a 3 diferenci na data $z = (1 : 10)^2$.

Řešení 5

```
1 > difference <- function(x, d = 1) {
2 +   # provedu diferenci
3 +   if (d > 0) {
4 +     x = (x[2:(length(x))] - x[1:(length(x) - 1)])
5 +   }
6 +   # pokud je difference vyssi nez jedna
7 +   # znovu difference funkce
8 +   if (d > 1) {
9 +     x = difference(x, d = d - 1)
10 +   }
11 +   return(x)
12 + }
13
14 > z = (1:10) ^ 2
15 > difference(x = z, d = 0)
16 [1] 1 4 9 16 25 36 49 64 81 100
17 > difference(x = z, d = 1)
18 [1] 3 5 7 9 11 13 15 17 19
19 > difference(x = z, d = 2)
20 [1] 2 2 2 2 2 2 2 2
21 > difference(x = z, d = 3)
22 [1] 0 0 0 0 0 0
```

Daná funkce provádí diferenci do té doby, než je hodnota parametru *d* rovna jedné – v tu chvíli provede poslední diferenci a navrátí výsledek. Funkce by šla zapsat ještě jednodušeji, a to pomocí jen jedné podmínky, jak ilustruje kód níže.

```
1 > difference <- function(x, d = 1) {
2 +   if (d > 1) {
3 +     x = (x[2:(length(x))] - x[1:(length(x) - 1)])
4 +     x = difference(x, d = d - 1)
5 +   } else if (d == 1) {
6 +     x = (x[2:(length(x))] - x[1:(length(x) - 1)])
7 +   }
8 +   return(x)
9 + }
```

První postup řešení pomocí dvou *if* podmínek zde uvádíme z důvodu přehlednosti a lepší čitelnosti.

Kapitola 6

Základy tvorby grafů v jazyce R

6.1 Úvodem ke grafům v R

Jednou z hlavních předností jazyka **R** jsou grafické možnosti vizualizace dat. V **R** lze vytvářet velmi kvalitní grafické výstupy, a to jak po stránce technické, tak i vizuální. Grafy v **R** představují samy o sobě velkou oblast, kde o každé zde prezentované kapitole vycházejí samostatné publikace. Z tohoto důvodu se omezíme na základy daných technologií. Obě prezentované knihovny jsou samostatným nezávislým systémem pro tvorbu grafů, který stojí osamoceně od ostatních prezentovaných knihoven. To je důležité si uvědomit zejména proto, že postupy, které lze aplikovat, například na graf knihovny `ggplot2`, nelze aplikovat na grafy z knihovny `base graphics` či jiné. Velkou výhodou těchto prezentovaných knihoven je i to, že existuje mnoho balíčků, které je rozšiřují a přidávají nové možnosti zpracování grafů (více viz. Rahlf, 2019; Venables, Smith, R Core Team, 2020). Vždy ale mějme na paměti, že nelze použít balíček na rozšíření jedné knihovny pro úpravu grafu z knihovny jiné.

Nelze říci, že jedna z prezentovaných knihoven na vizualizaci dat je jedinou vhodnou na používání a ostatní není třeba znát. Každá z těchto knihoven má výhody a nevýhody při běžném použití. Bez ohledu na použitou knihovnu (`ggplot`, `base graphic`, `lattice`) je základ grafiky vytvořen pomocí grafického engine `grDevices`. Tato knihovna je součástí výchozí instalace **R** a při každém zapnutí je automaticky zavedená. Základní grafy jazyka **R** (`Base graphics`) jsou postavené na knihovně `graphics` (grafický systém). Druhým grafickým systémem je pak systém `grid`, na kterém je postavená grafická knihovna `ggplot2` (nebo například i knihovna `lattice`). Z těchto důvodů jsou tyto knihovny navzájem nekompatibilní a nelze jejich výstupy ani nijak spojovat.

Velkým rozdílem je chování těchto knihoven z pohledu přístupu k objektům a použití funkcí. `Base graphics` nevytváří prioritně objekt, který obsahuje graf. Naopak knihovna `ggplot2` vytváří objekt, do kterého lze graf uložit a dále s ním i pracovat. S knihovnou

`ggplot2` se pracuje více po stránce objektového programování než funkcionálního. Knihovna `ggplot2` navíc pracuje oproti ostatním knihovnám ve vrstvách, tedy každý bod na grafu je tvořen samostatnou grafickou vrstvou.

Shrneme-li systém tvorby grafů v jazyku **R**, tak rozlišujeme:

- grafický engine (`grDevices`),
- grafické systémy (`grid` a `graphics`),
- grafické balíčky (`base graphics`, `lattice`, `ggplot2...`),
- balíčky rozšiřující grafické balíčky,
- balíčky rozšiřující grafický engine.

Dalším důležitým pojmem je tzv. nízkoúrovňová a vysokoúrovňová funkce. Vysokoúrovňové funkce jsou odpovědné za tvorbu grafů (vytváří grafické okno a jádro samotného grafu), zatímco nízkoúrovňové funkce jsou takové funkce, které jsou odpovědné za editaci grafu či jakoukoli následnou práci s ním.

6.2 Base graphics

Základní knihovnou pro tvorbu grafů je `base graphics` (R Core Team 2018, 2020). Jedná se o nativní funkce jazyka **R**, které jsou součástí výchozí instalace. Tyto grafy poskytují základní grafické nástroje pro vizualizaci. Jejich velkou výhodou je jednoduchost na ovládání a rychlost tvorby grafického výstupu. Nevýhodou je naopak nemožnost komplexního zobrazení dat mimo standardní postupy, popřípadě je to velmi složité (v porovnání například s `ggplot2`). Velkou nevýhodou je již zpětná, nemožná editace už vizualizovaných prvků grafu. Jakékoliv ex-post úpravy grafu se dějí překrytím/přetištěním. `Base graphics` má pro každý typ grafu specifickou funkci. Avšak různých výstupů lze dosáhnout i pomocí hlavní funkce (`plot()`), která podle typu objektu změní druh výstupu. Základní sada grafů v `base graphics` je následující:

- `plot()` – vytváří základní grafy, jako je liniový graf, scatter plot či lze použít i např. pro zobrazení map,
- `barplot()` – sloupcový graf,
- `boxplot()` – krabičkový diagram/graf,
- `hist()` – histogram,
- `dotchart()` – bodový graf neboli cleveland plot,
- `pie()` – výsečový (koláčový) graf,

-
- `stripchart()` – bodový graf pro vizualizaci vektorů dat.

R umožňuje v základu i více druhů grafů, kterými se zde nebudeme zabývat, ale jedná se například o grafy pro vizualizaci dvou a více proměnných jako jsou hvězdnicové/párovací diagramy (`stars()`), zobrazení obrázku (`image()`), zobrazení 3D prostoru (např. `contour()`, `persp()`) či mozaikové grafy (`mosaicplot()`) anebo grafy pro vizualizaci změn distribucí, četností a jiného určení (jako například `cdplot()`, `spineplot()`, `fourfoldplot()`).

Důležité je mít na paměti to, že mnohé funkce se chovají jinak při jiné třídě datového vstupu. Tuto problematiku si ukážeme následně u jednotlivých typů grafů.

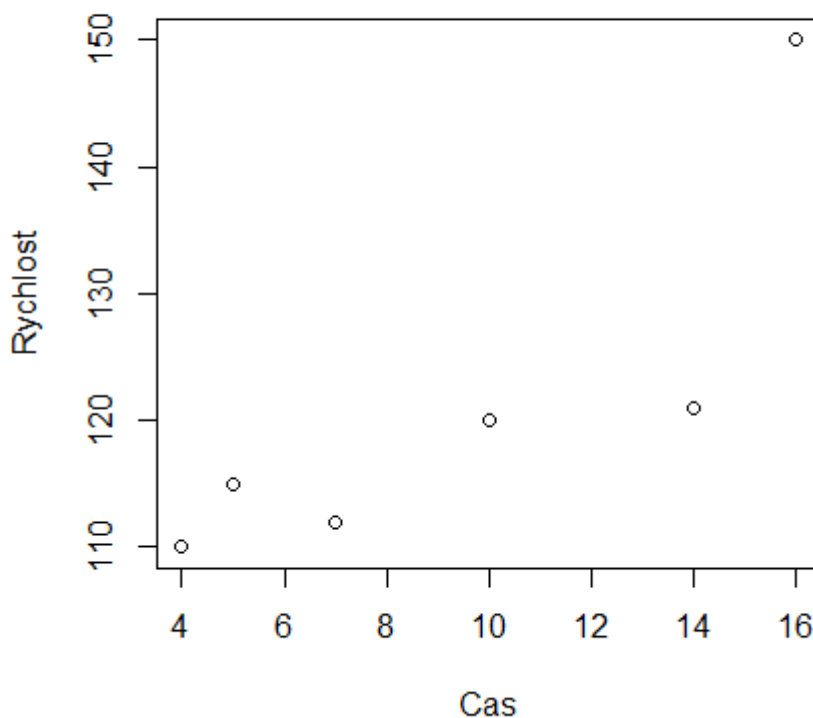
6.2.1 Funkce `plot()`

Vytvořme si prvně datový soubor, který chceme vizualizovat. Bude se jednat o data popisující rychlost a čas blíže neurčeného vozu.

```
1 > Auto <- data.frame(Cas = c(4, 5, 7, 10, 14, 16),  
2 +                      Rychlost = c(110, 115, 112, 120, 121, 150))
```

Pokud do funkce `plot()` vložíme objekt typu `data frame` o dvou sloupcích, tak vykreslí bodový graf těchto proměnných. Na ose x bude uvedená první proměnná a na ose y bude uvedená druhá. **R** dále automaticky přebere názvy těchto proměnných a vloží je jako popisek jednotlivých os. Automaticky se též změní („přeškálují“) osy grafu podle rozsahu daných proměnných. Graf lze vizualizovat pomocí příkazu `plot()`.


```
1 > plot(Auto)
```



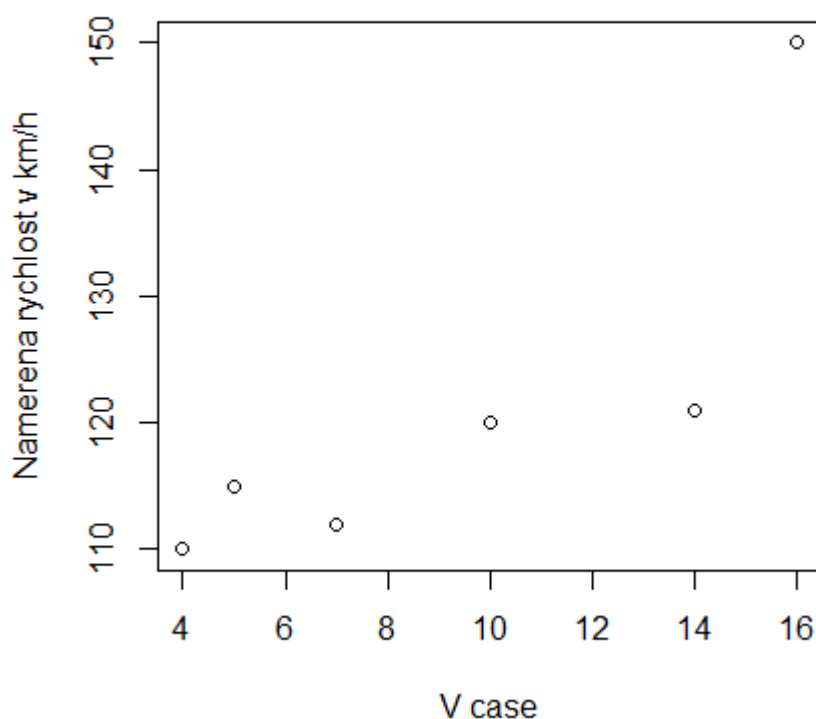
Obrázek 6.1: Funkce plot()

Zdroj: Vlastní zpracování.

Graf lze vizualizovat i pomocí příkazu `plot(x = Auto$Cas, y = Auto$Rychlost)`. Avšak nevýhodou tohoto zápisu je to, že místo rozumného názvu os se nám vloží text volání dané proměnné z data frame (tedy `Auto$Cas` a `Auto$Rychlost`). Zde je pak nutné všechny ostatní popisky vložit pomocí proměnných funkce `plot`.

```
1 > plot(Auto, main = "Rychlost auta v závislosti na case",  
2 +     sub = "Zdroj: Vlastní zpracování.", xlab = "V case",  
3 +     ylab = "Namerena rychlost v km/h")
```

Rychlost auta v závislosti na case



Zdroj: Vlastní zpracování.

Obrázek 6.2: Funkce `plot()` – úprava popisků

Na předcházejícím příkladu je pak vidět, že lze i oba přístupy kombinovat. Sice jsme předali data ve formátu data frame, ale řekli jsme jazyku R, jaké hodnoty má vykreslit do grafu. Parametr `main` určuje hlavní nadpis grafu, `sub` pak podnadpis. Popisky os určují parametry `xlab` a `ylab`. U většiny grafů mají jednotlivé parametry analogické použití. Až na situace, jako jsou parametry pro editaci atributů grafu, který daný graf nemá – například osy u koláčového grafu aj. To, jaké máme základní atributy, si ukážeme hned v následující podkapitole.

Formule

Většina grafů podporuje krom standardního zápisu, kdy přiřazujeme parametrům funkce odpovídající proměnné i zápis pomocí vzorce. Vzorec je zde chápán jako typ proměnné – formule. Ve formuli obvykle používáme základní aritmetická znaménka a místo rovná se symbol vlnovky `~`. Toto znaménko má využití zejména v regresi, kdy potřebujeme

zapsat rovnici, kterou odhadujeme. Kdybychom použili znaménko rovná se (=), nedošlo by k vytvoření objektu formule, ale ke snaze o přiřazení uprostřed funkce. V případě grafů slouží znaménko tilda (vlnovka) k rozlišení toho, co má být vůči čemu vykreslené. V některých aplikacích se můžeme setkat i s tečkou, tečka zastupuje všechny ostatní proměnné datasetu. Máme tedy tři způsoby, jak vytvořit identický graf, který se ale bude lišit ve výchozích popiscích os.

```
1 > plot(Auto)
2 > plot(x = Auto$Cas, y = Auto$Rychlost)
3 > plot(Rychlost ~ Cas, data = Auto)
```

6.2.2 Základní editace grafu a barvy

Parametry ovlivňující tvorbu grafů můžeme rozdělit do dvou skupin. Ty, které ovlivňují jak vysokoúrovňové funkce, tak i nízkoúrovňové, a ty, které lze použít pouze v nízkoúrovňových funkcích (pomocí editace funkcí par). Následující tabulka shrnuje základní vysokoúrovňové grafické parametry funkcí jazyka **R**.

Parametr	Popis
x, y	parametr pro předání hodnot jednotlivých os
type	druh grafu
cex	velikost písma (v násobkách)
main, sub	hlavní nadpis a podnadpis
xlab, ylab	popisky os
fg	barva osy grafu
lwd	tloušťka linie
lty	druh linie
col	barva linie či jiných vizualizovaných prvků v oblasti grafu
pch	druh symbolu, kterým se mají data vizualizovat
family, font	font a typ písma
xlim, ylim	rozsah os grafu
las	natočení/rotace popisků os

Tabulka 6.1: Parametry funkcí

Zdroj: Vlastní zpracování.

Parametr **cex** ovlivňuje relativní velikost objektů ve vztahu k velikosti řádky. Barvu os můžeme ve funkci **plot** změnit pomocí **fg** parametru. Zde je ale důležité poznamenat, že **R** dokáže přebrat různé definice barev. Základní barvy lze získat pomocí řadových čísel

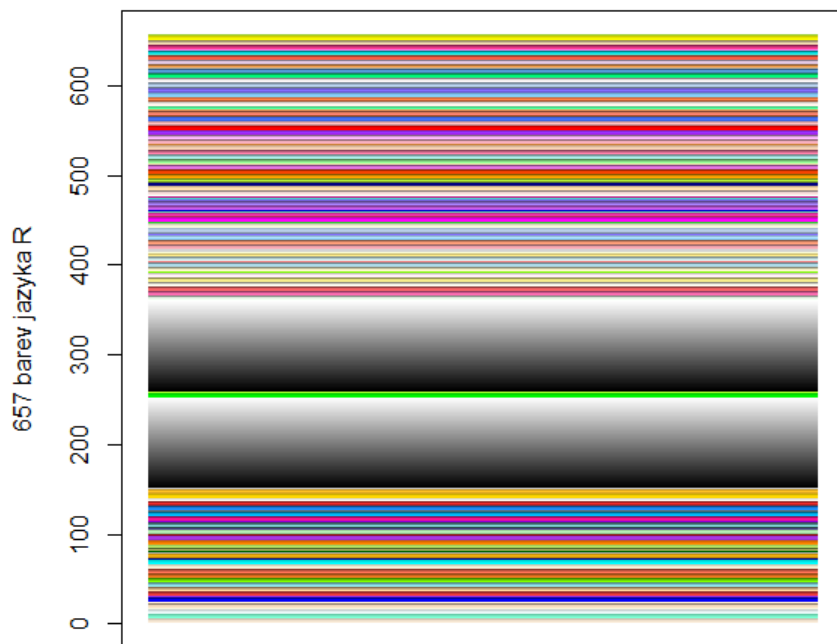
0–8. Dále lze využít i anglická jména barev. Seznam základních jmen 657 barev lze získat pomocí příkazu `colors()`.

```
1 > colors()
2 [1] "white" "aliceblue" "antiquewhite"
3 [4] "antiquewhite1" "antiquewhite2" "antiquewhite3"
4 [7] "antiquewhite4" "aquamarine" "aquamarine1"
5 [10] "aquamarine2" "aquamarine3" "aquamarine4"
6 [13] "azure" "azure1" "azure2"
7 [16] "azure3" "azure4" "beige"
8 [19] "bisque" "bisque1" "bisque2"
9 [22] "bisque3" "bisque4" "black"
10 [25] "blanchedalmond" "blue" "blue1"
11 [28] "blue2" "blue3" "blue4"
12 ...
13 ...
```

Jak je vidět z kódu výše, R má pro každý typ základních 657 barev své jméno. Obvykle jsou dané barvy seskupené podle odstínu a valéru. Pomocí následujícího kódu si můžeme dané barvy i zobrazit. To, jak následující kód funguje a co jaké parametry způsobují, si představíme postupně v rámci této kapitoly.

```
1 > plot(
2 +   y = c(1:657),
3 +   x = c(1:657),
4 +   main = "Barvy v R",
5 +   type = "n",
6 +   xlab = "",
7 +   ylab = "657 barev jazyka R",
8 +   sub = "Zdroj: vlastní práce autora",
9 +   xaxt = 'n'
10 + )
11 > for (i in 1:657) {
12 +   lines(
13 +     x = c(1:657),
14 +     y = c(rep(i, 657)),
15 +     col = colors()[i],
16 +     lwd = 0.01
17 +   )
18 + }
```

Barvy v R



Zdroj: vlastní zpracování.

Obrázek 6.3: Základních 657 barev jazyka R

Jazyk **R** ale podporuje i `rgb`, `hsv` a `hlc` reprezentaci barev. Tyto reprezentace lze použít pomocí analogických funkcí: `rgb()`, `hsv()` a `hlc()`.

To, jakou grafickou formou se budou data vizualizovat, ovlivňuje parametr `type`. Zde máme na výběr mezi více druhy reprezentace dat:

p	body
l	linie (čára)
b	kombinace bodů a linie
c	jako b, ale s vynecháním bodů
o	vykreslí body a pak přes ně vykreslí linii
h	jako histogram
s	schody
n	nic

Tabulka 6.2: Druhy čar

Zdroj: Vlastní zpracování.

Některé z těchto parametrů (`cex`, `col`, `font`) mají případné varianty pro editaci specifických parametrů (`main`, `sub`, `axis`, `lab`) a editují se pomocí spojení těchto dvou názvů s tečkou, například `cex.main` bude ovlivňovat pouze relativní velikost nadpisu, `col.lab` jen barvu popisků os. Fonty, které si můžeme vybrat pomocí parametru `family`, jsou: „sans“, „serif“, „mono“ a „symbol“. Tuto sadu fontů lze samozřejmě rozšířit pomocí dodatečných balíčků. Parametr `font` pak specifikuje druh písma: 1 – plain, 2 – bold, 3 – italic, 4 – bold a italic, 5 – symbol. Parametry `xlim` a `ylim` se zadávají jako atomický numerický vektor o délce dva, kde první hodnota označuje minimum a druhá hodnota maximum osy grafu. V následujícím obrázku zkusíme použít větší množství těchto parametrů.

```

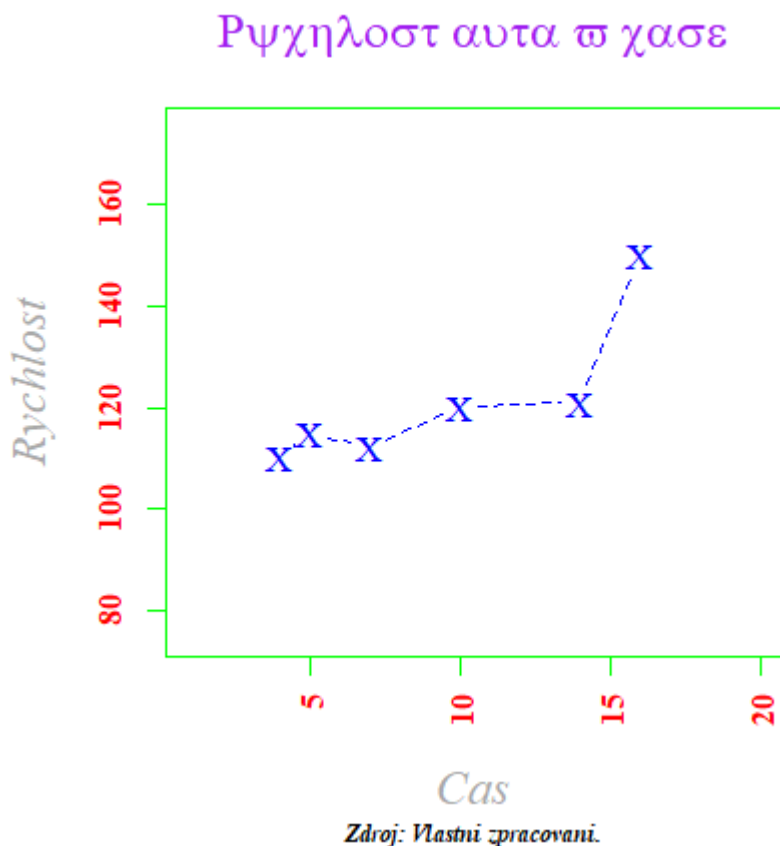
1 > plot(
2 +   Auto,
3 +   type = "b",
4 +   lty = "dashed",
5 +   pch = "X",
6 +   cex = 1.2,
7 +   cex.main = 1.5,
8 +   cex.sub = 0.75,
9 +   cex.lab = 1.5,
10 +  family = "serif",
11 +  font.main = 5,
12 +  font.sub = 4,
13 +  font.lab = 3,
14 +  font.axis = 2,
15 +  col = "blue",
16 +  col.axis = "red",
17 +  col.lab = "darkgrey",
18 +  col.main = "purple",
19 +  fg = "green",
20 +  main = "Rychlost auta v case",

```

```

21 + sub = "Zdroj: Vlastni zpracovani",
22 + xlim = c(1, 20),
23 + ylim = c(75, 175),
24 + las = 3
25 + )

```

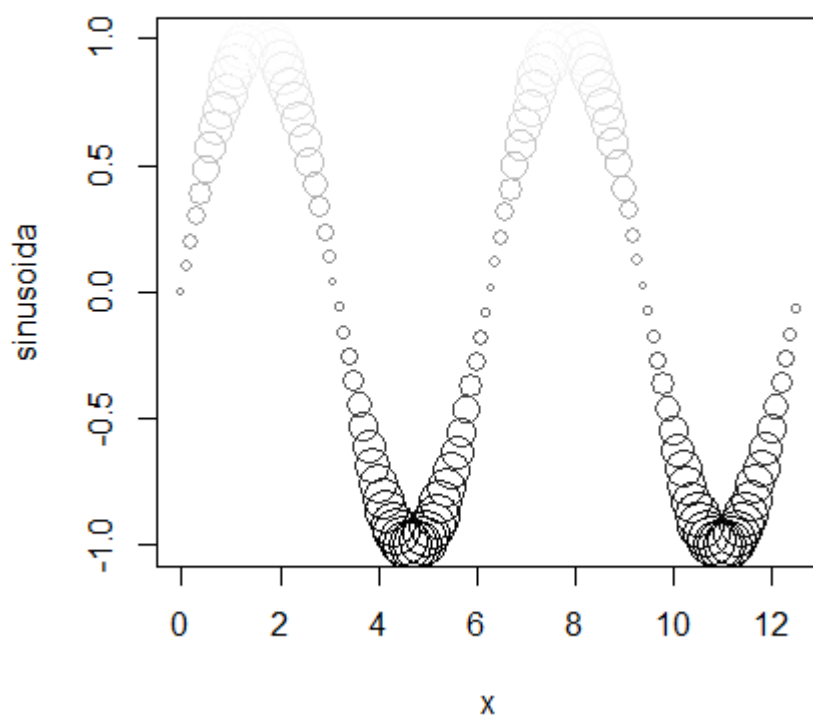


Obrázek 6.4: Funkce plot() – úprava parametrů

R obsahuje celou škálu palet barev, které lze využít při tvorbě grafů. Jedná se například o funkce `rainbow()`, `hlc.colors()`, `heat.colors()`, `terrain.colors()` či `topo.colors()`. Tyto funkce přebírají jako hlavní parametr počet barev, které se z dané palety mají vygenerovat (v rgb kódu). Alternativou je pak například funkce `grey()`, která pro hodnoty mezi 0 a 1 přiřadí odpovídající barvu na škále šedé. Použití funkce `grey()` je výrazně jednodušší, například v následujícím kódu si vytvoříme sekvenci hodnot od 0 do čtyřnásobku π s krokem 0,1 (obrázek 6.4). Vypočítáme si hodnoty sinusoidy pomocí funkce `sin` a vykreslíme graf. V tomto grafu budeme chtít, aby: 1) hodnoty byly na škále šedi (tmavší hodnoty pod nulou a světlejší nad) a 2) aby hodnoty, které jsou vzdálenější od 0, měly

větší symbol. Toho lze snadno dosáhnout tak, že do parametru `col` přiřadíme ne jednu hodnotu, ale počet hodnot odpovídajících jednotlivým pozorováním. Pak se každá daná hodnota vykreslí požadovaným atributem. Funkce `grey` pro hodnoty blízké 0 vykresluje černou barvu a pro hodnoty blízké jedničce bílou. Je tedy nutné vypočítané hodnoty sinusoidy vhodně přeskálovat. Toto uděláme tak, že k její hodnotě přičteme jedna a vydělíme dvěma. Tímto dosáhneme požadované škály šedi. V případě velikosti nám pak stačí zvolit výchozí (nejmenší) velikost bodu (zde jsme zvolili 0,5) a pak zvolit vhodné škálování. My jsme absolutní hodnotu sinusoidy násobili třemi, aby byla změna velikosti velmi patrná.

```
1 > x <- seq(from = 0, to = pi * 4, by = 0.1)
2 >
3 > sinusoida <- sin(x)
4 >
5 > plot(sinusoida ~ x,
6 +     col = grey(c(sinusoida + 1) / 2),
7 +     cex = 0.5 + abs(sinusoida) * 3)
```

Obrázek 6.5: Graf sinusoidy

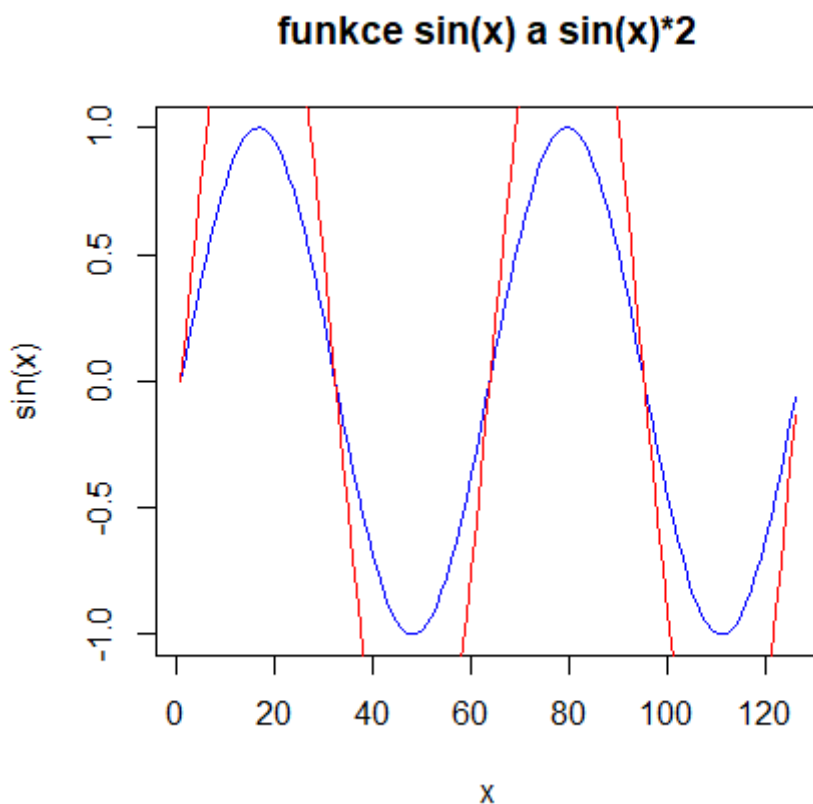
Zdroj: Vlastní zpracování.

6.2.3 Základní typy grafů

Funkce `plot()` – více proměnných

Než si představíme všechny ostatní typy grafů, tak se ještě vrátíme k základní funkci `plot()`. Ukázali jsme si, jak vykreslit jednoduchý graf, ale co když chceme vykreslit více proměnných v rámci jednoho grafu? Máme dvě možnosti. Buď použijeme funkci `points()`, anebo `lines()`. Obě tyto funkce použijeme až poté, co vykreslíme samotný graf. A obě fungují velmi podobně jako funkce `plot()`. Avšak s tím, že neinicilizují, oproti funkci `plot()`, novou grafickou oblast, ale překreslují stávající. S tímto chováním se ale pojí jeden zásadní problém. Pokud jsme již něco „vytiskli“ v rámci funkce `plot()`, pak to již nejsme schopni opravit funkcí `lines` či `points`. Často se pak uživatelům stává, že vykreslí jednu proměnnou, a když chtějí vykreslit druhou, tak zjistí, že se nevejde do grafické oblasti grafu a jaksi „mizí“ mimo rámec grafu (nebo se rovnou vůbec nezobrazí, protože se tyto dvě proměnné ani úrovněově neprotínají). Tomuto lze předejít vhodným nastavením parametrů `xlim` a `ylim` v rámci funkce `plot()`. Tedy když plánujeme použít funkci `lines` či `points`, je nutné již s tímto počítat při funkci `plot`. Ukažme si toto na funkci `sin(x)`, `2*sin(x)`. Následující příklad ukazuje, co se stane, pokud nenastavíme vhodně proměnné `xlim` a `ylim`.

```
1 > # Tvorba promene
2 > x = seq(from = 0, to = pi * 4, by = 0.1)
3 > # Vypocet sinusoidy
4 > sinusoida = sin(x)
5 > # Graficka funkce plot
6 > plot(sinusoida,
7 +     type = "l",
8 +     col = "blue",
9 +     xlab = "x",
10 +     main = "funkce sin(x) a sin(x)*2"
11 + )
12 > # Pomoci funkce lines
13 > # pridame dalsi caru do grafu
14 > lines(sin(x) * 2,
15 +     type = "l",
16 +     col = "red")
```

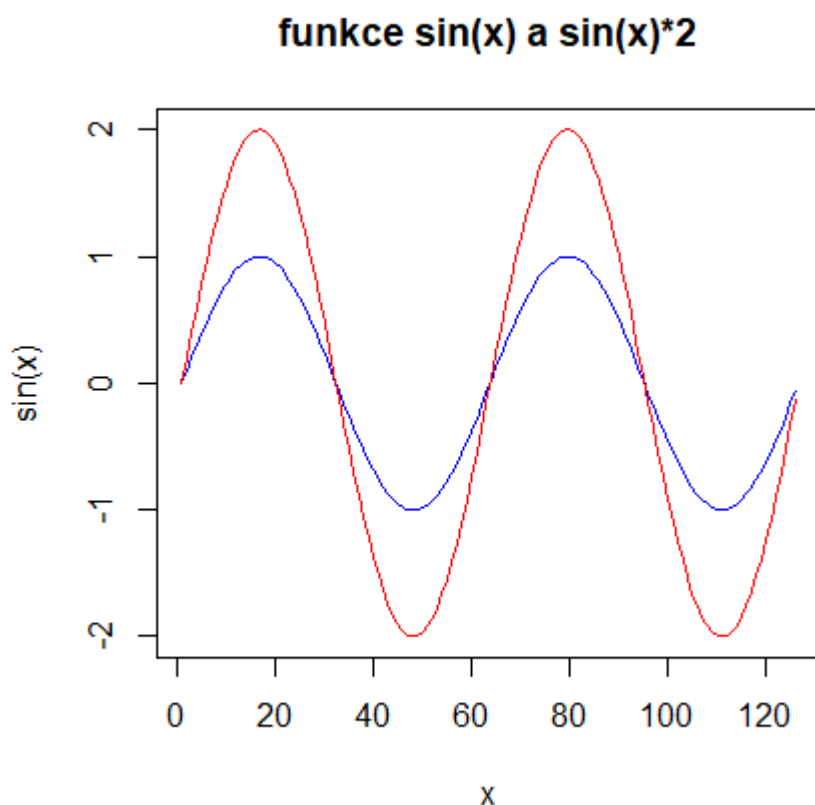


Obrázek 6.6: Funkce `plot()` a více proměnných

Zdroj: Vlastní zpracování.

Poměrně snadno lze toto vyřešit pomocí parametru `ylim`. Najdeme minimální a maximální hodnotu v obou proměnných a tu předáme do požadovaného parametru.

```
1 > x = seq(from = 0, to = pi * 4, by = 0.1)
2 >
3 > plot(sin(x), type = "l", col = "blue", xlab = "x",
4 +     main = "funkce sin(x) a sin(x) * 2",
5 +     ylim = c(min(c(sin(x), sin(x) * 2)),
6 +             max(c(sin(x), sin(x) * 2))))
7 >
8 > lines(sin(x) * 2, type = "l", col = "red")
```



Obrázek 6.7: Funkce `plot()` – `ylim` a `xlim` parametry

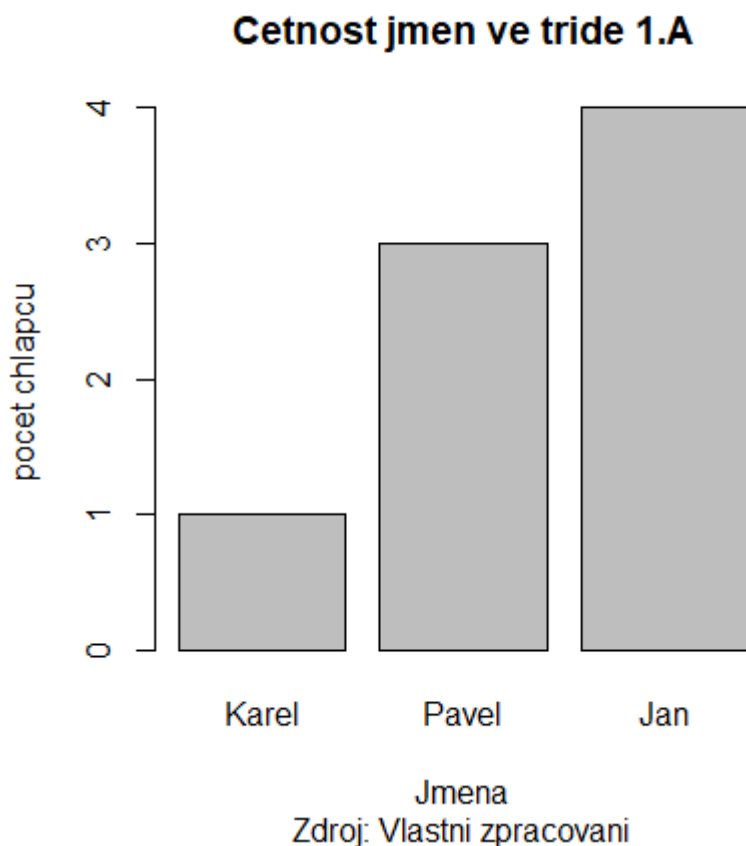
Zdroj: Vlastní zpracování.

Funkce `barplot()`

Sloupcový graf neboli anglicky „barplot“ je typ grafu, kde jsou data rozdělená do určitých skupin (obvykle na ose x) a vizualizována absolutní (či relativní) četností (nebo součtem) v dané skupině, kterou lze vyčíst na ose y . Oproti příkazu `plot` nepřebírá tato funkce parametr `x` či `y`, ale parametr `height`, kterému mají odpovídat výšky daných sloupců. Výšku daného sloupce (četnost či jiný údaj) si musíme spočítat sami. Funkce `barplot` dokáže přebrat mnoho typů objektů tak jako funkce `plot`. Standardní postup je vložit do této funkce pojmenovaný atomický numerický vektor.

```
1 > n <- c(1, 3, 4)
2 > names(n) <- c("Karel", "Pavel", "Jan")
3 > barplot(n, main = "Četnost jmen ve třídě 1.A",
4 +       sub = "Zdroj: Vlastní zpracování",
```

```
5 + ylab = "pocet chlapcu", xlab = "Jmena")
```



Obrázek 6.8: Funkce `barplot()`

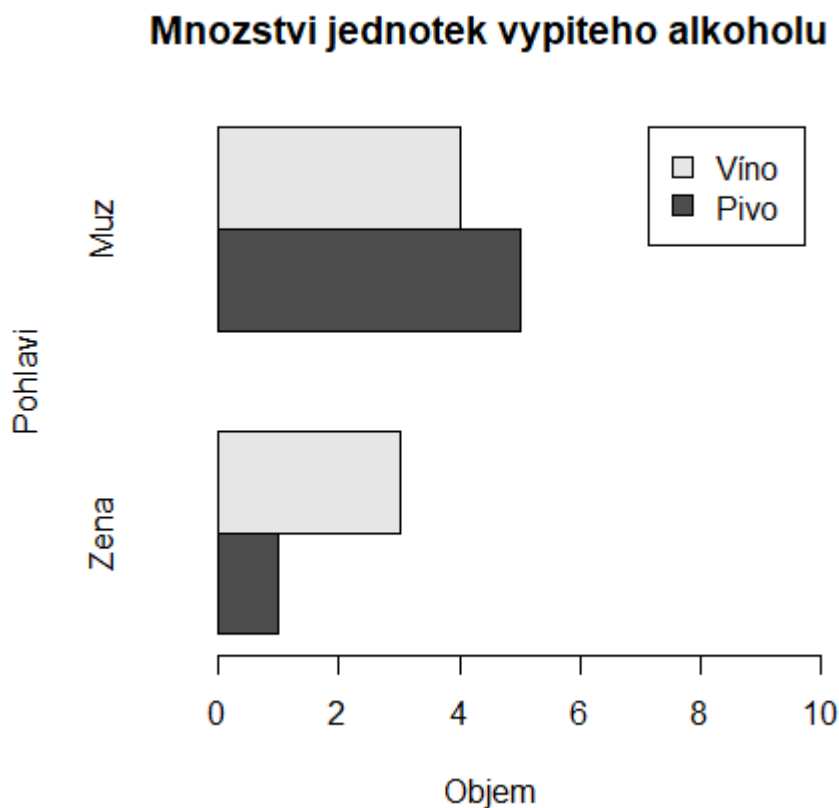
Jak může pozorný čtenář vidět, funkce `barplot` má mnoho společných parametrů, které lze použít k editaci grafu. Tato funkce má navíc mnoho dalších vlastních parametrů – např. `space` ovlivňuje množství místa mezi sloupci, `width` ovlivňuje šířku sloupce (musí být použit ale zároveň i parametr `xlim`) do parametru `names.arg` (funguje i parametr `namesn`, ale ostatní nastavení nemusí fungovat zcela správně) můžeme specifikovat alternativní jména sloupců, `beside` v případě zobrazení dat více kategorií pro každou skupinu (z matice) umožní volit typ vizualizace dat podkategorií (vedle sebe/na sobě). Parametr `col` přebírá barvu, popř. barvy sloupců.

```
1 > Alkohol <- matrix(c(1, 3, 5, 4), nrow = 2,
2 +                   dimnames = list(c("Pivo", "Vino"),
3 +                                   c("Zena", "Muz")))
4 > barplot(Alkohol, beside = TRUE, horiz = TRUE, xlim = c(-1, 10),
```

```

5 +     main = "Mnozstvi jednotek vypiteho alkoholu",
6 +     ylab = "Pohlavi", xlab = "Objem",
7 +     legend.text = c("Pivo", "Vino"))

```



Obrázek 6.9: Úprava parametrů funkce `barplot()`

Zdroj: Vlastní zpracování.

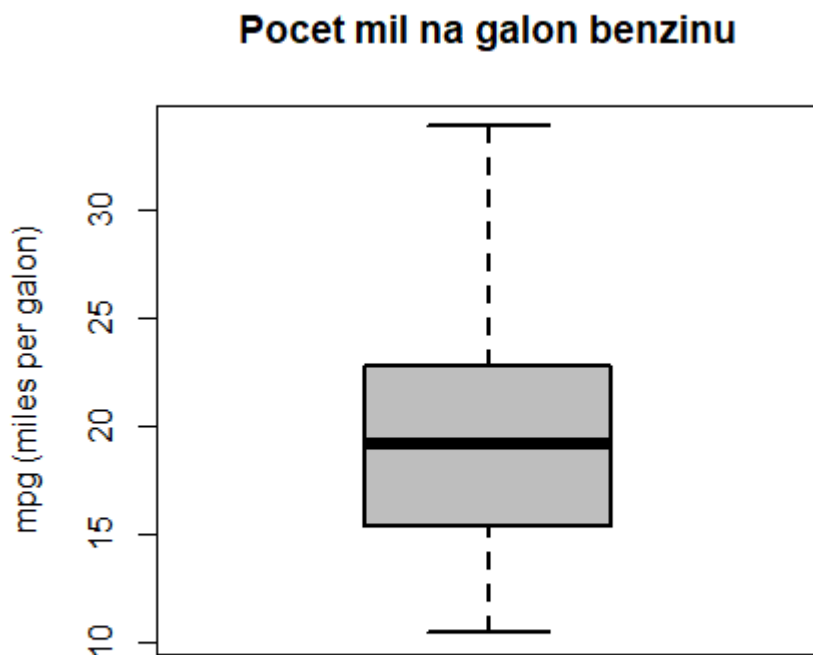
V případě, kdy chceme vytvořit horizontální orientaci grafu, pak parametru `horiz` přiřadíme hodnotu `TRUE` a graf se nám otočí.

Funkce `boxplot()`

Krabičkový diagram vytvoříme snadno pomocí funkce `boxplot()`. Funkce `boxplot()` umožňuje i vytváření více krabičkových diagramů. Toho lze dosáhnout buď předáním objektu typu `data.frame`, anebo vzorcem, který nám řekne, jak má data dělit. Funkce

`boxplot()` přebírá primárně data ve formátu `data.frame`, anebo jakýkoliv vektor (atomický vektor nebo list). Ukážeme si použití tohoto grafu na datovém souboru `mtcars`. Nejprve si zkusíme vykreslit krabičkový diagram počtu mil, které ujedou jednotlivá auta na galon benzínu (proměnná `mpg`).

```
1 > boxplot(  
2 +   x = mtcars$mpg,  
3 +   main = "Pocet mil na galon benzinu",  
4 +   sub = "Dataset mtcars",  
5 +   ylab = "mpg (miles per gallon)",  
6 +   lwd = 2,  
7 +   col = "grey"  
8 + )
```

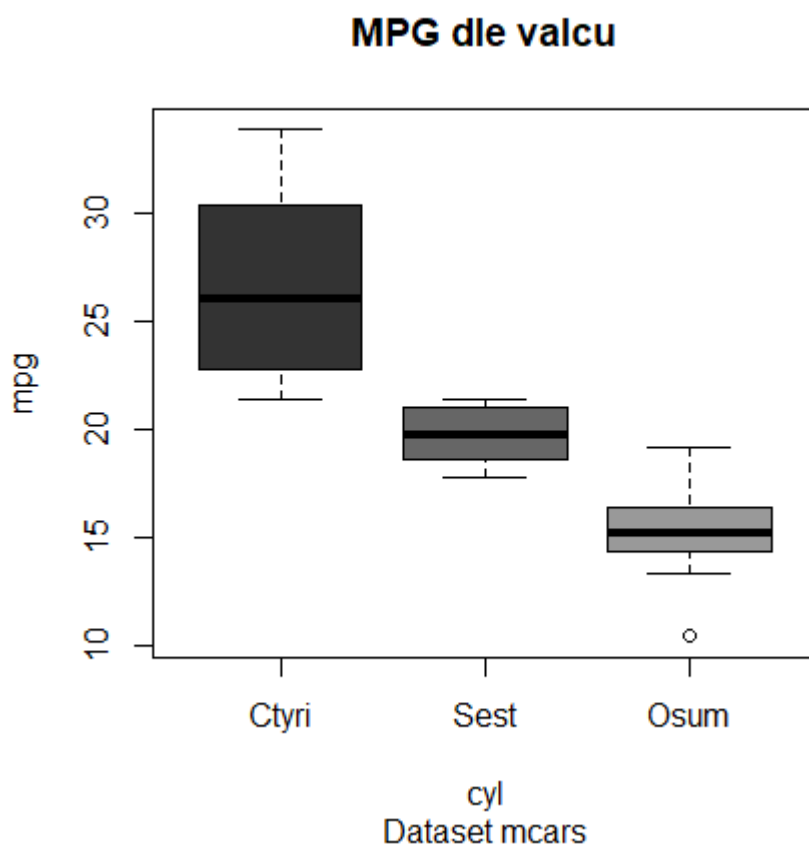


Obrázek 6.10: Funkce `boxplot()`

Zdroj: Vlastní zpracování.

Funkce `boxplot()` umí efektivně vykreslovat i více krabičkových diagramů v jedné grafické oblasti. Řekněme, že si budeme chtít vykreslit krabičkové diagramy toho, kolik mil ujede auto na jeden galon benzínu (`mpg`) pro jednotlivé varianty počtů válců auta (`cyl`). Máme dvě možnosti, jak takovýto graf vytvořit. Buď vyrobíme nový data frame objekt, který bude obsahovat data o tom, kolik ujedou daná auta při jednotlivých počtech válců, nebo využijeme nám již známé formule.

```
1 > boxplot(  
2 +   mpg ~ cyl,  
3 +   data = mtcars,  
4 +   col = grey(c(0.2, 0.4, 0.6)),  
5 +   names = c("Ctyri", "Sest", "Osum"),  
6 +   lwd = 1.5,  
7 +   main = "MPG dle valcu",  
8 +   sub = "Dataset mcars"  
9 + )
```

Obrázek 6.11: Funkce `boxplot()` – úprava parametrů funkce

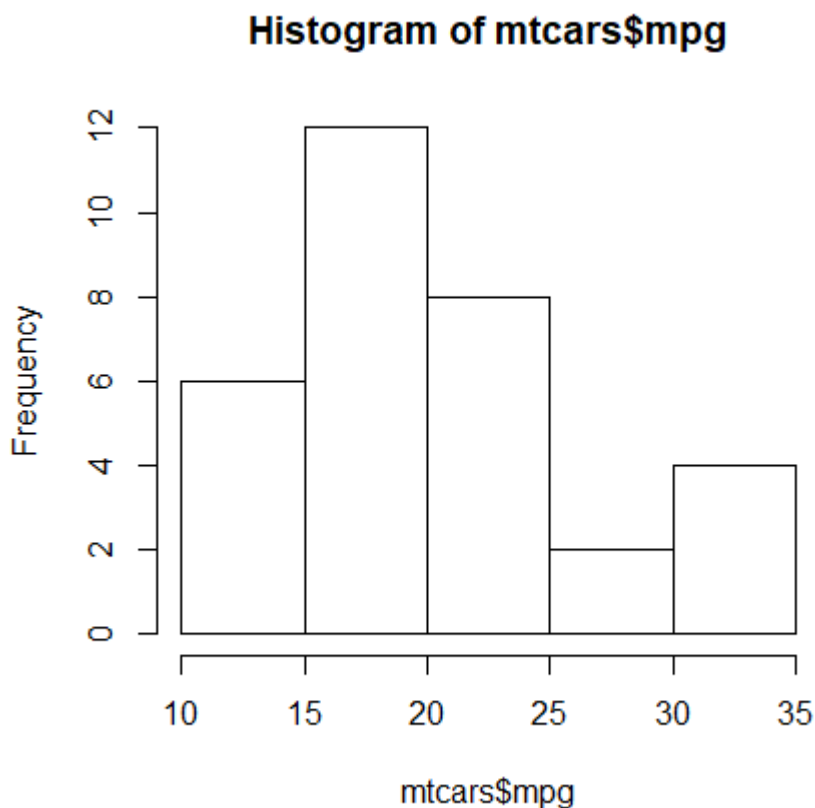
Zdroj: Vlastní zpracování.

Jak vidíme, i v případě boxplotu lze nastavit individuální parametry pro dílčí krabičkové diagramy. Zde jsme pomocí vektoru na škále šedi nastavili každému grafu jinou barvu. Jména na ose x se řídí parametrem `names` (ne `names.arg` jako v předchozím případě).

Funkce `hist()`

V **R** můžeme histogram vytvořit velmi snadno a rychle pouze jedním příkazem – `hist()`. **R** se za vás bude snažit spoustu věcí vyřešit samo - jako například zvolit počet sloupců histogramu. Pojdme si nejprve ukázat, jak vytvořit snadno histogram, a pak si jej zkusíme upravit.

```
1 hist(mtcars$mpg)
```

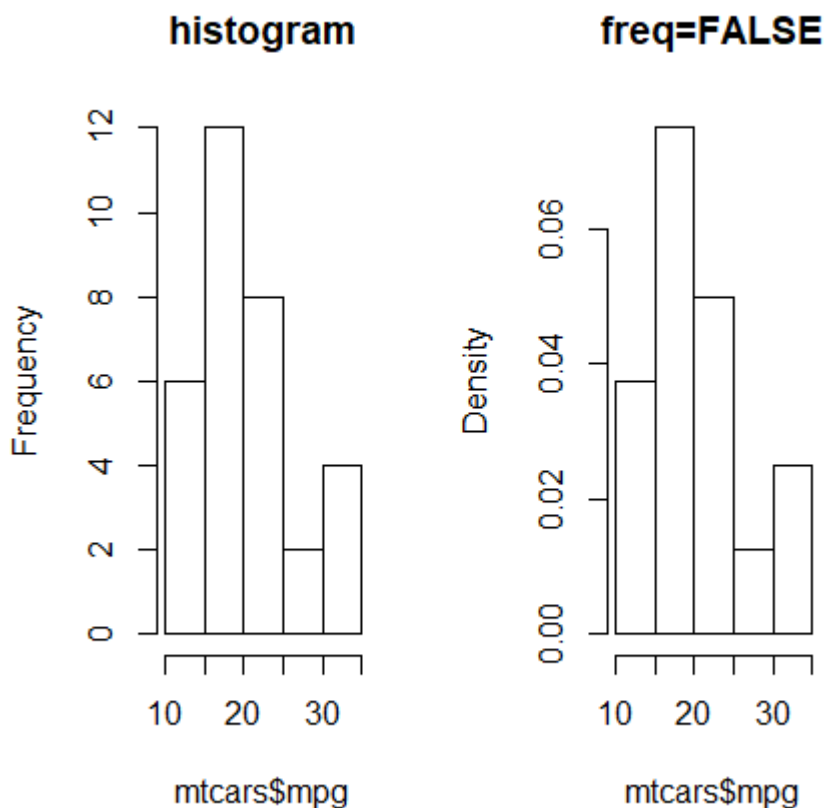


Obrázek 6.12: Funkce hist()

Zdroj: Vlastní zpracování.

Z příkladu vidíme, že **R** za nás volí automaticky nadpis histogramu. Tak jako v případě funkce plot ale můžeme nadpis a popisky os editovat pomocí parametrů `main`, `xlab` a `ylab`. Histogram, který **R** ve výchozím nastavení vykreslí, obsahuje na ose y absolutní četnosti. Pokud bychom chtěli vytvořit graf hustoty, stačí vložit do parametru `freq` hodnotu `FALSE`. Na následujícím obrázku si porovnáme oba grafy.

```
1 > par(mfrow = c(1, 2))
2 > hist(mtcars$mpg, main = "histogram")
3 > hist(mtcars$mpg, freq = FALSE, main = "freq=FALSE")
```



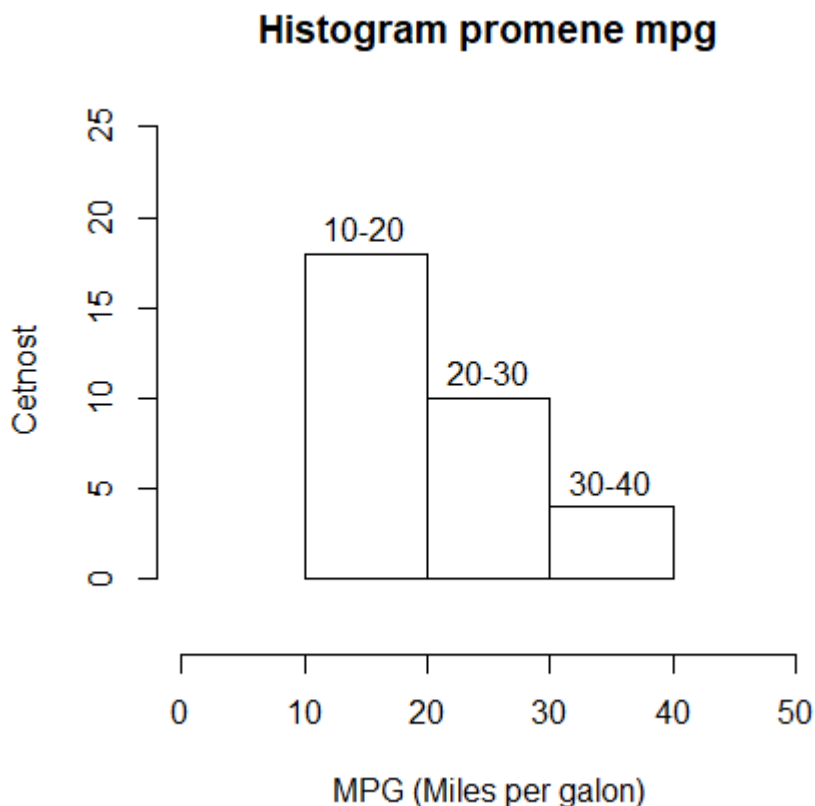
Obrázek 6.13: Funkce `hist()` a `hist(,freq=FALSE)`

Zdroj: Vlastní zpracování.

Pozorný čtenář si jistě všimne, že oba grafy jsou „tvarem“ sloupců identické. Rozdíl mezi „density“ histogramem a „frequency“ histogramem spočívá v ose y – tedy v tom, v jakých jednotkách je výsledný graf reprezentován. Počet sloupců můžeme ovlivnit pomocí parametru `breaks`. Tento parametr je schopen ovlivnit širokou škálu vstupů. Pokud mu předáme jedno číslo, tak ho použije na počet intervalů, které má vytvořit, pokud mu předáme vektor několika čísel, pak tyto hodnoty bude považovat za hraniční hodnoty jednotlivých intervalů sloupců histogramu (podobně jako u funkce `cut()`). Dále je schopen tento parametr přebrat i funkci, která mu řekne, jak má vytvořit intervaly pro sloupce histogramu. V neposlední řadě přebírá jméno matematické metody, jež se má použít na výpočet počtu sloupců a histogramu. Tento parametr je při tvorbě histogramů klíčový a ovlivní výslednou vypovídající hodnotu dat. Dále je důležitý v situaci, kdy chceme mezi sebou porovnat jednotlivé histogramy. Pokud bychom měli data do sloupců rozdělena jiným způsobem, bylo by porovnání velmi obtížné. Dalším důležitým parametrem je `freq`.

Tento parametr definuje, zda se mají data vizualizovat absolutní četností. Parametry `xlim`, `ylim` fungují zcela identicky jako v případě funkce `plot()`.

```
1 > hist(  
2 +   mtcars$mpg,  
3 +   breaks = c(10, 20, 30, 40),  
4 +   labels = c("10-20", "20-30", "30-40"),  
5 +   main = "Histogram promenne mpg",  
6 +   ylim = c(-3, 25),  
7 +   xlim = c(0, 50),  
8 +   xlab = "MPG (Miles per gallon)",  
9 +   ylab = "Cetnost"  
10 + )
```



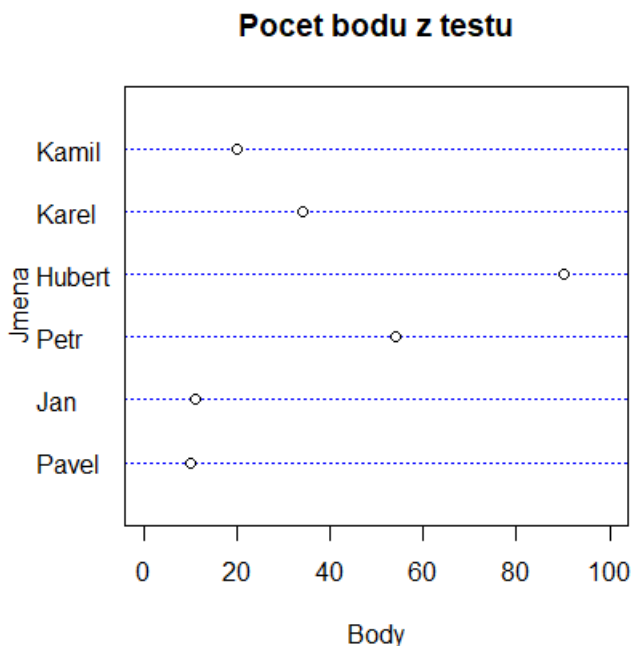
Obrázek 6.14: Funkce `hist()` – úprava parametrů funkce

Zdroj: Vlastní zpracování.

Funkce dotchart()

Bodový diagram (Cleveland bodový graf) vytvoříme pomocí funkce `dotchart()`. Tato funkce přebírá data, která vizualizuje do parametru `x`. Funkce umí zpracovat jak vektor, tak matice. Pokud funkce přebere jako hlavní parametr vektor, tak vizualizuje jeden bodový diagram. V případě, že přebere matici, tak **R** vytvoří více samostatných bodových diagramů v rámci jednoho grafu – rozdělený do skupin podle sloupců matice. Následující příklad znázorňuje počet bodů, které získali jednotliví studenti v rámci školního turnaje.

```
1 > body <- c(10, 11, 54, 90, 34, 20)
2 > names(body) <- c("Pavel", "Jan", "Petr", "Hubert", "Karel", "Kamil")
3 >
4 > dotchart(body, main = "Pocet bodu z testu",
5 +           xlim = c(0, 100), lcolor = "blue",
6 +           xlab = "Body", ylab = "Jmena")
7 + )
```



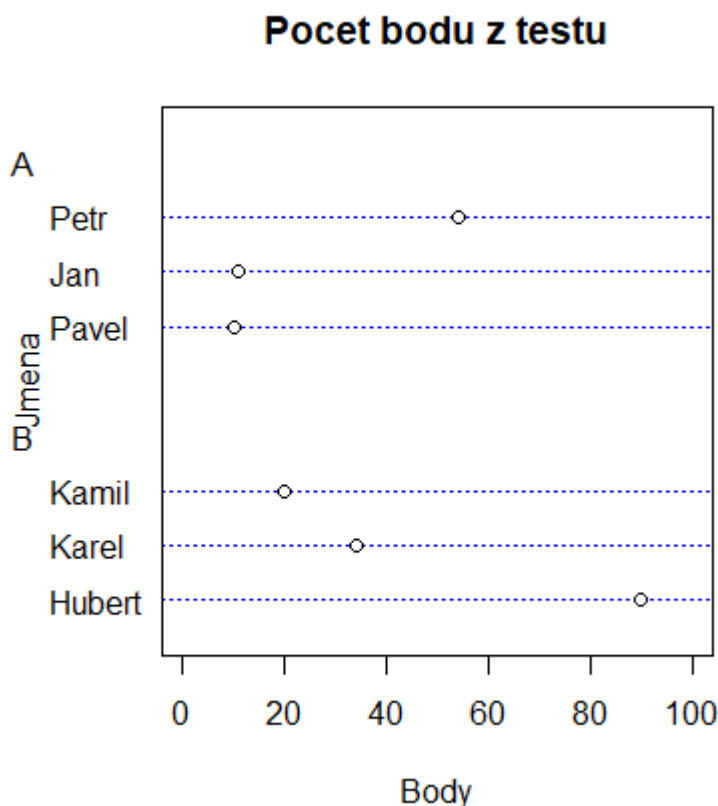
Obrázek 6.15: Funkce `dotchart()`

Zdroj: Vlastní zpracování.

Pokud bychom chtěli daná data rozdělit do skupin dle jejich třídy, tak bychom si museli vytvořit nový vektor (nejlépe proměnné typu faktor), který nám pomůže identifikovat, kdo

do jaké třídy patří. Výhodou tohoto postupu je, že nám zabezpečí snadnou srovnatelnost obou grafů.

```
1 > body <- c(10, 11, 54, 90, 34, 20)
2 > names(body) <- c("Pavel", "Jan", "Petr", "Hubert", "Karel", "Kamil")
3 > tridy <- as.factor(c("A", "A", "A", "B", "B", "B"))
4 > dotchart(
5 +   body,
6 +   main = "Pocet bodu z testu",
7 +   xlim = c(0, 100),
8 +   lcolor = "blue",
9 +   xlab = "Body",
10 +  ylab = "Jmena",
11 +  groups = tridy
12 + )
```



Obrázek 6.16: Funkce `dotchart()` – úprava parametrů funkce

Zdroj: Vlastní zpracování.

V případě, kdy by proměnná třídy (`tridy`) nebyla faktorem (a byla numerickou proměnnou!), budou data sice opticky rozdělena do skupin, ale bez popisku těchto skupin. V tomto případě je popiskem pouze A a B, ale lze vytvořit faktorovou proměnnou podle své potřeby.

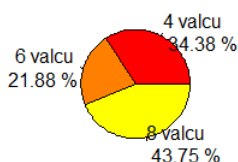
Funkce `pie()`

Výsečový graf znázorňuje data poměrným zastoupením výšece kruhu na celém jeho obsahu. Jedná se o velmi populární graf, avšak musíme si uvědomit, že tento typ grafu může výsledky opticky zkreslovat – zejména, když nejsou k dispozici popisky výsečí s jejich hodnotou. Proto si vytvoříme graf, který bude obsahovat jak popisky, tak i procentní zastoupení jednotlivých výsečí na celku. V následujícím příkladu si budeme vizualizovat zastoupení aut v data setu dle počtu válců motoru auta. Pro rychlé shrnutí nám postačí

použít funkci `table`.

```
1 > p.cyl <- table(mtcars$cyl) / length(mtcars$cyl) * 100
2 > pie(
3 +   x = p.cyl,
4 +   main = "Procentni zastoupeni zastoupeni aut \n dle poctu
      cylindru",
5 +   sub = "Dataset mtcars",
6 +   labels = paste(names(p.cyl), "valcu \n", round(p.cyl, 2), "%")
7 +   ,
8 +   col = heat.colors(3)
9 + )
```

**Procentni zastoupeni zastoupeni aut
dle poctu cylindru**



Dataset mtcars

Obrázek 6.17: Funkce `pie()`

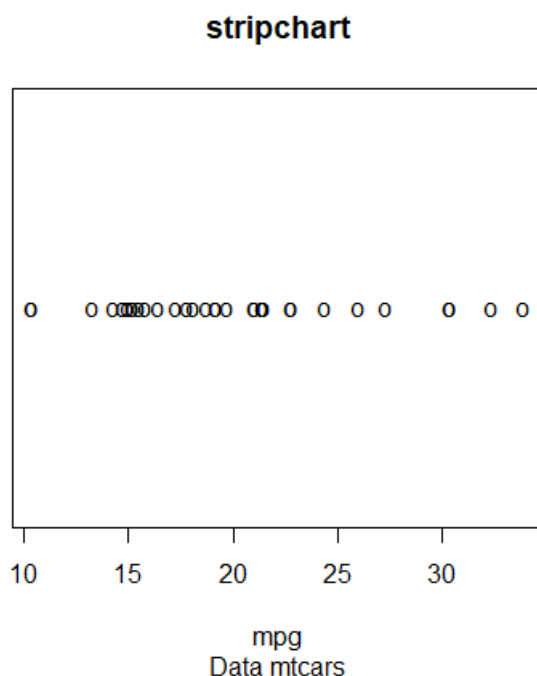
Zdroj: Vlastní zpracování.

Pozorný čtenář si jistě všimne, že zde pro vykreslení barev není již použita funkce `grey()`, ale `heat.colors()`. Funkce `heat.colors()` navrátí přechod „ohnivých“ barev při zadaném počtu barev. Barvy zde neslouží k identifikaci toho, zda má jedna výšeč větší hodnotu, ale k jejich odlišení. Stejně fungují funkce `grey.colors()`, `topo.colors()`, `terrain.colors()`, `cm.colors()` nebo `rainbow()`. Další sady barev můžeme získat například prostřednictvím knihovny `ColorBrewer`. Tato knihovna disponuje širokou škálou palet pro různá použití. Například zde najdeme jak gradientní přechody barev, tak ale i divergenční škály. Divergenční škály jsou takové, kdy máme plynulý přechod z jedné barvy do druhé. Například z modré do červené.

Funkce `stripchart()`

Funkce `stripchart()` vytváří v základu pro uživatele možná na první pohled velmi zvláštní graf, ale pokud si dáme práci a nastavíme jednotlivé parametry této funkce správně, můžeme získat velmi efektní graf s velkou informační hodnotou. Stripchart je vlastně 1-dimenzionální verze scatter plotu – verze bodového grafu jen s jednou proměnnou. Ukažme si nejprve výchozí graf.

```
1 > stripchart(x = mtcars$mpg, main = "stripchart",  
2 +   xlab = "mpg", sub = "Data mtcars", pch = "o")
```



Obrázek 6.18: Funkce `stripchart()`

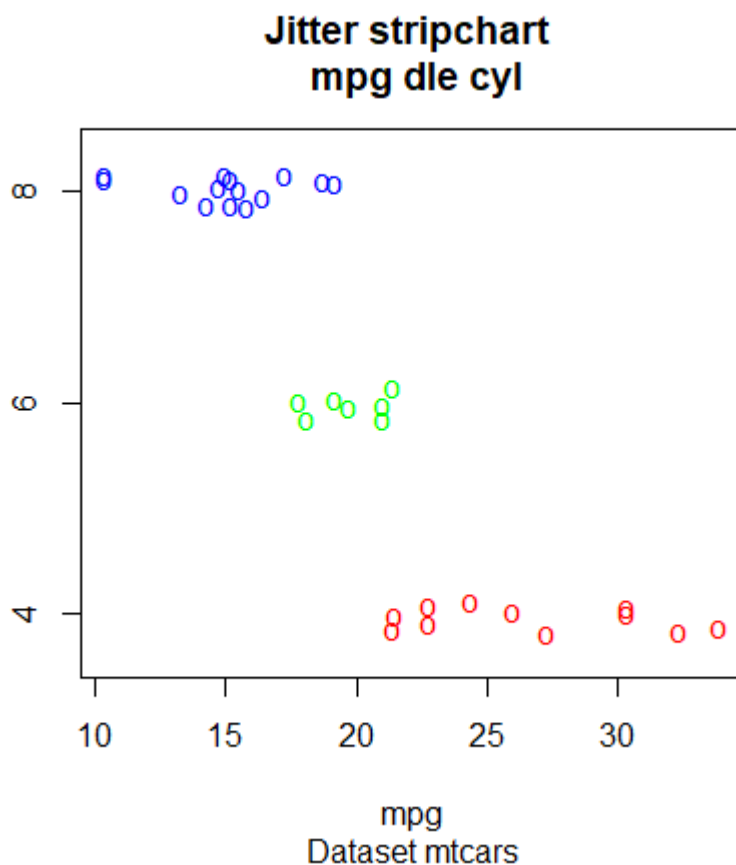
Zdroj: Vlastní zpracování.

Na první pohled se může jednat o nepřehledný graf. Máme zde jen jednu osu – x. Na této ose jsou vizualizované hodnoty z předaného vektoru `mpg`. Z grafu je vidět, že se mnoho hodnot navíc překrývá, a pokud by nabývaly přibližně stejné hodnoty, tak nejsme schopni opticky posoudit „zahuštění“ bodů v dané části grafu. Důležitým parametrem je zde `jitter`. Tento parametr ovlivňuje to, jak jsou body zahuštěné v místě jejich překrytí. Pokud nastavíme metodu z `overplot` na `jitter`, tak se aktivuje vertikální rozhození bodů. Další výhodou tohoto grafu je, že umí přebrat vzorec, a tak jsme schopni rozdělit data například pomocí proměnné `mpg` a podle počtu válců (cyl proměnná).

```

1 > stripchart(mtcars$mpg ~ mtcars$cyl,
2 +   xlab = "mpg", pch = "o", lwd = 1.5, jitter = 0.1,
3 +   method = "jitter", col = rainbow(3),
4 +   main = "Jitter stripchart \n mpg dle cyl",
5 +   sub = "Dataset mtcars"
6 + )

```



Obrázek 6.19: Funkce stripchart() se specifikací jitter

Zdroj: Vlastní zpracování.

Na tomto příkladu je vidět, že nastavení parametrů funkce pro tvorbu grafů je velmi důležité a ovlivňuje, jak čtenář graf vnímá a jakou informační hodnotu pro něj má. Z prvního grafu vytvořeného pomocí funkce `stripchart()` nešlo mnoho údajů vyčíst. Naopak z druhého je snadno patrné významné rozdělení spotřeby benzínu v závislosti na počtu válců.

6.2.4 Funkce `par()` a ostatní nízkoúrovňové funkce

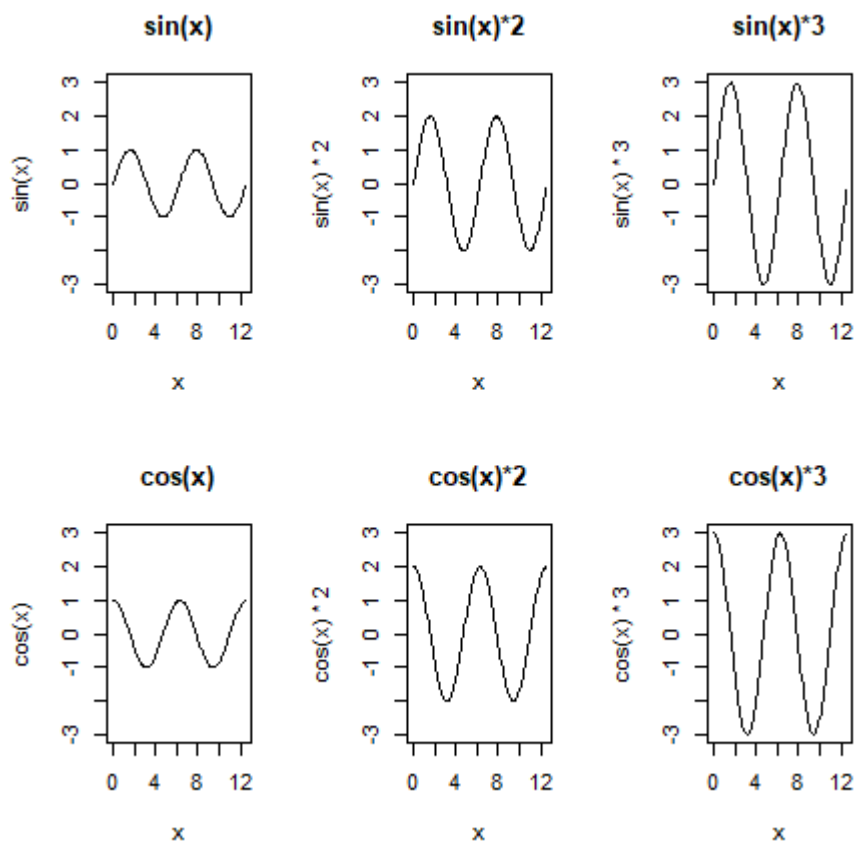
Na začátku kapitoly jsme si představili parametry funkcí, které jsou použitelné zejména pro vysokoúrovňové grafické funkce (tedy funkce tvořící graf). Nedílnou součástí jazyka **R** jsou ale i nízkoúrovňové funkce. Z nich se teď podrobně podíváme na tu nejdůležitější – funkci `par()`. Funkce `par` je funkce, která, když se zavolá, tak navrátí vektor typu `list`, na jehož prvcích jsou klíčová grafická nastavení pro tvorbu grafu. Nízkoúrovňová funkce nikdy neovlivňuje graf, který je již zobrazen, popřípadě grafické okno, které je otevřeno. Zde je důležité upozornit na významný rozdíl mezi výchozím prostředím **R** a RStudiem.

Více grafů v jedné tiskové oblasti

V případě, kdy chceme do jedné tiskové oblasti vytvořit více grafů, máme na výběr mezi dvěma možnostmi. Buď chceme vytvořit grafické okno, kde každý graf bude mít stejné zastoupení z pohledu prostoru a tyto grafy budou srovnané v řádkách/sloupcích (parametry `mfrow/mfcol` z `par` funkce), anebo musíme využít `layout` funkce. Třetí možností by bylo použít funkci `split.screen()`. Avšak této poslední možnosti je lepší se vyvarovat. Vytváří mnohé problémy s editací a následným ukládáním grafů.

Nízkoúrovňová funkce `par` má dva parametry, které umožňují vložit více grafů do jedné grafické oblasti. Parametr `mfrow` přebírá vektor, kde je na prvním místě počet řádků grafů a na druhém místě počet sloupců grafů, které má budoucí grafická oblast obsahovat. Parametr `mfcol` má pořadí prvků ve vektoru, který přebírá jen přehozené – na prvním místě počet sloupců a na druhém počet řádků grafů. Nastavením jednoho parametru automaticky funkce `par()` přenastaví ten druhý. Uživatel se nikdy nesmí snažit oba parametry nastavit najednou. Například budu-li si chtít vykreslit funkce $\sin(x)$, $\cos(x)$ a jejich násobky dvěma a třemi v jednom grafickém okně, ale každý graf zvlášť. Výsledkem budou tedy tři grafy vedle sebe a dva řádky, jeden pro funkci sinus a jeden pro funkci cosinus.

```
1 > x = seq(from = 0, to = pi * 4, by = 0.1)
2 > par(mfrow = c(2, 3))
3 > plot(sin(x) ~ x, main = "sin(x)",
4 +     type = "l", ylim = c(-3, 3))
5 > plot(sin(x) * 2 ~ x, main = "sin(x)*2",
6 +     type = "l", ylim = c(-3, 3))
7 > plot(sin(x) * 3 ~ x, main = "sin(x)*3",
8 +     type = "l", ylim = c(-3, 3))
9 > plot(cos(x) ~ x, main = "cos(x)",
10 +    type = "l", ylim = c(-3, 3))
11 > plot(cos(x) * 2 ~ x, main = "cos(x)*2",
12 +    type = "l", ylim = c(-3, 3))
13 > plot(cos(x) * 3 ~ x, main = "cos(x)*3",
14 +    type = "l", ylim = c(-3, 3))
```

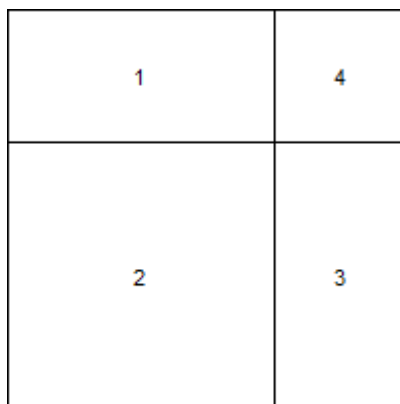


Obrázek 6.20: Více grafů v jedné grafické oblasti – funkce `par()`

Zdroj: Vlastní zpracování.

Další důležitý rozdíl mezi parametrem `mfrow` a `mfcol` spočívá v pořadí, v jakém jsou jednotlivé oblasti pro grafy naplňovány. Parametr `mfrow` postupuje po řádcích. **R** začíná vyplňovat v levém horním rohu a postupuje po jednotlivých řádcích. Naopak parametr `mfcol` způsobí, že grafy jsou umísťované do buněk po sloupcích. Pro vyplňování grafů lze použít i cykly. Alternativou k tomuto postupu je vytvoření si vlastního grafického modelu pomocí funkce `layout()`. Funkce `layout` přebírá jako hlavní parametr matici, kde jednotlivé buňky matice představují grafy. Buňky této matice můžeme graficky spojit tak, že jim přiřadíme stejnou hodnotu. Zároveň pořadí čísel definuje pořadí naplňování grafické oblasti grafy.

```
1 m <- matrix(c(1, 2, 2, 1, 2, 2, 4, 3, 3), ncol = 3)
2 my.layout <- layout(m)
3 layout.show(my.layout)
```

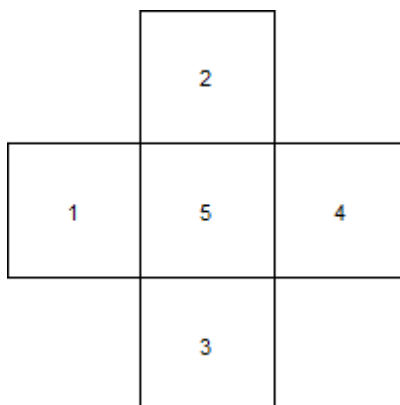


Obrázek 6.21: Layout model grafu

Zdroj: Vlastní zpracování.

Funkce `layout.show()` nám pak umožní vytvořit graf našeho grafického modelu. Tedy funkce `layout().show` nevytváří budoucí graf, ale ukazuje nám, jaký grafický model se snažíme vytvořit. Dalším důležitým parametrem funkce `layout` je `width` a `height` – šířka a výška. Tyto dva parametry nám pomohou ovlivnit šířku jednotlivých sloupců a řádků layout modelu. Parametr `n` pak definuje celkový počet grafů, kterým naplníme grafickou oblast.

```
1 > m <- matrix(c(0, 1, 1, 0, 2, 5, 5, 3, 2, 5, 5, 3, 0, 4, 4, 0),
  nrow = 4)
2 > my.layout <- layout(m, height = c(1, 2, 2, 1), width = c(1, 2, 2,
  1))
3 layout.show(my.layout)
```

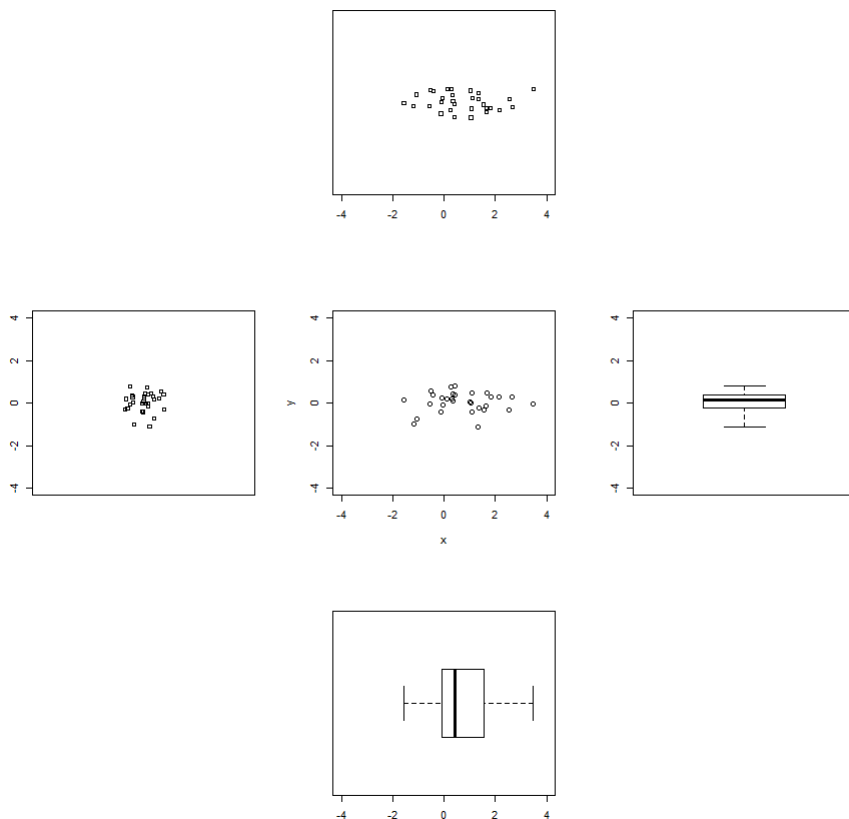


Obrázek 6.22: Layout model grafu II

Zdroj: Vlastní zpracování.

Na tomto grafickém modelu je zejména zajímavé použití nuly. Nula označuje oblasti, které nemají být zaplněny grafy. Oblast grafů nemusí být vůbec souvislá. Zkusíme si vytvořit pět grafů uspořádaných do kříže. Budeme popisovat dvě náhodně generované proměnné, kdy nalevo a nahoře budou jednodimenzionální bodové grafy a napravo a dole budou krabíčkové diagramy. Uprostřed se bude nacházet standardní bodový graf obou proměnných.

```
1 > set.seed(1)
2 > m <- matrix(c(0, 1, 1, 0, 2, 5, 5, 3, 2, 5, 5, 3, 0, 4, 4, 0),
3             nrow = 4)
3 > my.layout <- layout(m, height = c(2, 1, 1, 2), width = c(2, 1, 1,
4             2))
4 > y <- rnorm(n = 30, mean = 0, sd = 0.5)
5 > x <- rnorm(n = 30, mean = 0.5, sd = 1.5)
6 > stripchart(y,
7 +           method = "jitter",
8 +           vertical = TRUE,
9 +           ylim = c(-4, 4))
10 > stripchart(x, method = "jitter", xlim = c(-4, 4))
11 > boxplot(x, ylim = c(-4, 4), horizontal = TRUE)
12 > boxplot(y, ylim = c(-4, 4))
13 > plot(x, y, xlim = c(-4, 4), ylim = c(-4, 4))
```



Obrázek 6.23: Více grafů v jedné grafické oblasti – layout funkce

Zdroj: Vlastní zpracování.

V případě, když vám graf nejde vytvořit, je to nejspíš tím, že máte malou grafickou oblast. Pak stačí zvětšit grafickou oblast (okno, ve kterém se má graf vizualizovat). Nevyřešeným problémem u grafu 6.23 jsou okraje. Možná jste si všimli, že mezi grafy samotnými jsou příliš široké mezery. Na tento problém se podíváme hned v následující kapitole.

Nastavení okrajů grafu

Grafy v **R** mají dva druhy okrajů – vnitřní a vnější. Vnitřní okraj nastavujeme pomocí parametru **mar** (násobky řádků) či **mai** (v palcích). Vnější okraje nastavujeme pomocí **oma** (v řádkách) či **omi** (v palcích). Ve výchozím nastavení jsou vnější okraje nulové a vnitřní mají hodnotu parametru **mar**: 5.1, 4.1, 4.1 a 2.1. Tyto údaje jsou zaznamenány formou atomického numerického vektoru. První hodnota je spodní okraj, druhá levý okraj, třetí hodnota je horní okraj a poslední hodnota je pravý okraj. Hodnoty zapisujeme ve

směru hodinových ručiček od šesté hodiny. Dále je vhodné znát výchozí nastavení grafu v **R**. Řádky pod osou se číslují od nuly. Na nulté řádce jsou značky na osách, na první řádce pod osou jsou umístěné hodnoty příslušící ke značkám os. Další řádka se nechává standardně prázdná a na třetí řádce pak jsou popisky os (**xlab**/**ylab**). S ohledem na to, že se řádky číslují od nuly, je nutné mít okraje minimálně nastavené na hodnotu 4 u os (první a druhý okraj). V následujícím příkladu si na tyto jednotlivé řádky okolo grafu umístíme popisky. Popisky mimo **sub** a **main** lze vytvořit pomocí funkce **mtext()**. Do funkce se zadává text popisku (parametr **text**), strana (parametr **side**), řádka, na kterou se má text napsat (parametr **line**), a zda jde o vnější anebo vnitřní okraj (**outer**).

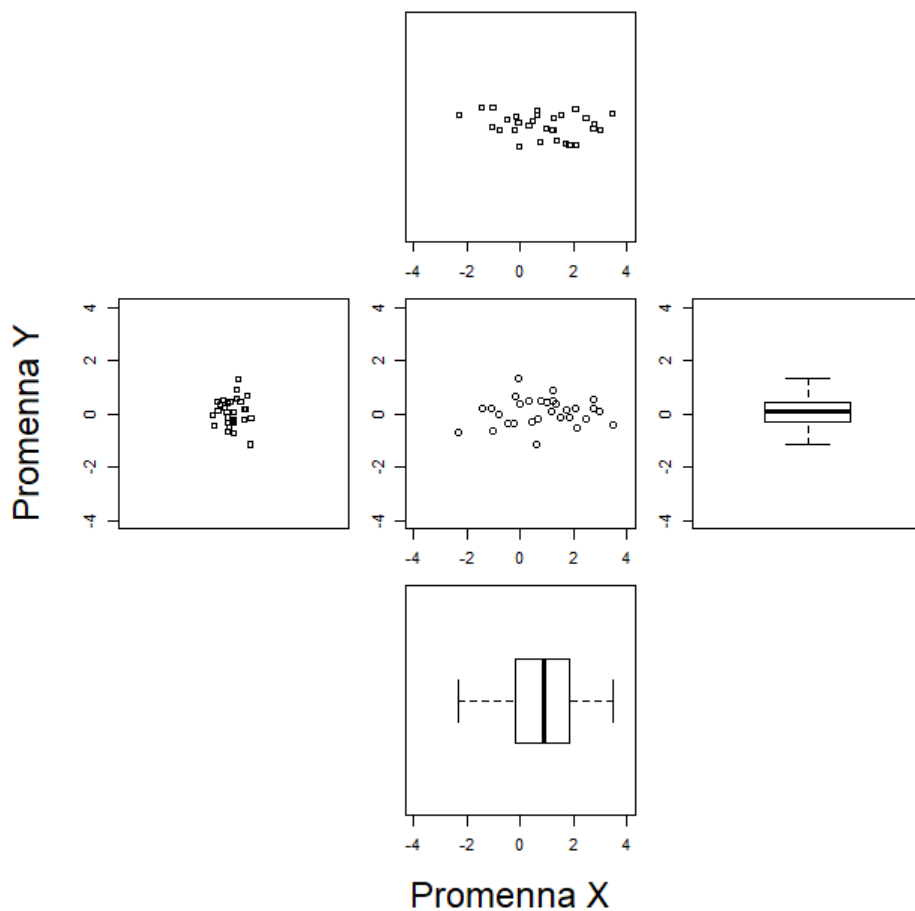
```
1 > # Funkci par nastavíme okraje
2 > par(mar = rep(4.1, 4), oma = rep(3, 4))
3 >
4 > # Základní ramec grafu
5 > plot(
6 +   c(0, 1),
7 +   type = "n",
8 +   xlab = "X - lab.....",
9 +   ylab = "Y - lab.....",
10 +   main = "Nadpis....."
11 + )
12 >
13 > # Inner okraj
14 > for (i in -2:4) {
15 +   mtext(
16 +     text = paste(i, "radka inner"),
17 +     side = 1:4,
18 +     line = i,
19 +     outer = FALSE
20 +   )
21 + }
22 > # Outer okraj
23 > for (i in 1:3) {
24 +   mtext(
25 +     text = paste(i, "radka outer"),
26 +     side = 1:4,
27 +     line = i,
28 +     outer = TRUE
29 +   )
30 + }
```



```

1 > set.seed(1)
2 >
3 > m <- matrix(c(0, 1, 1, 0, 2, 5, 5, 3, 2, 5, 5, 3, 0, 4, 4, 0),
4 +           nrow = 4)
5 >
6 > my.layout <- layout(m, height = c(2, 1, 1, 2), width = c(2, 1, 1,
7   2))
8 >
9 > y <- rnorm(n = 30, mean = 0, sd = 0.5)
10 >
11 > x <- rnorm(n = 30, mean = 0.5, sd = 1.5)
12 >
13 > # Nastaveni par() funkce
14 > par(mar = c(2, 2, 1, 1), oma = rep(4, 4))
15 >
16 > stripchart(y,
17 +           method = "jitter",
18 +           vertical = TRUE,
19 +           ylim = c(-4, 4))
20 >
21 > stripchart(x, method = "jitter", xlim = c(-4, 4))
22 >
23 > boxplot(x, ylim = c(-4, 4), horizontal = TRUE)
24 >
25 > boxplot(y, ylim = c(-4, 4))
26 >
27 > plot(x, y, xlim = c(-4, 4), ylim = c(-4, 4))
28 >
29 > # Spolecne popisky
30 > mtext(
31 +   text = "Promenna X",
32 +   side = 1,
33 +   line = 2,
34 +   outer = TRUE,
35 +   cex = 1.5
36 + )
37 >
38 > mtext(
39 +   text = "Promenna Y",
40 +   side = 2,
41 +   line = 2,
42 +   outer = TRUE,
43 +   cex = 1.5
44 + )

```



Obrázek 6.25: Úprava okrajů grafu

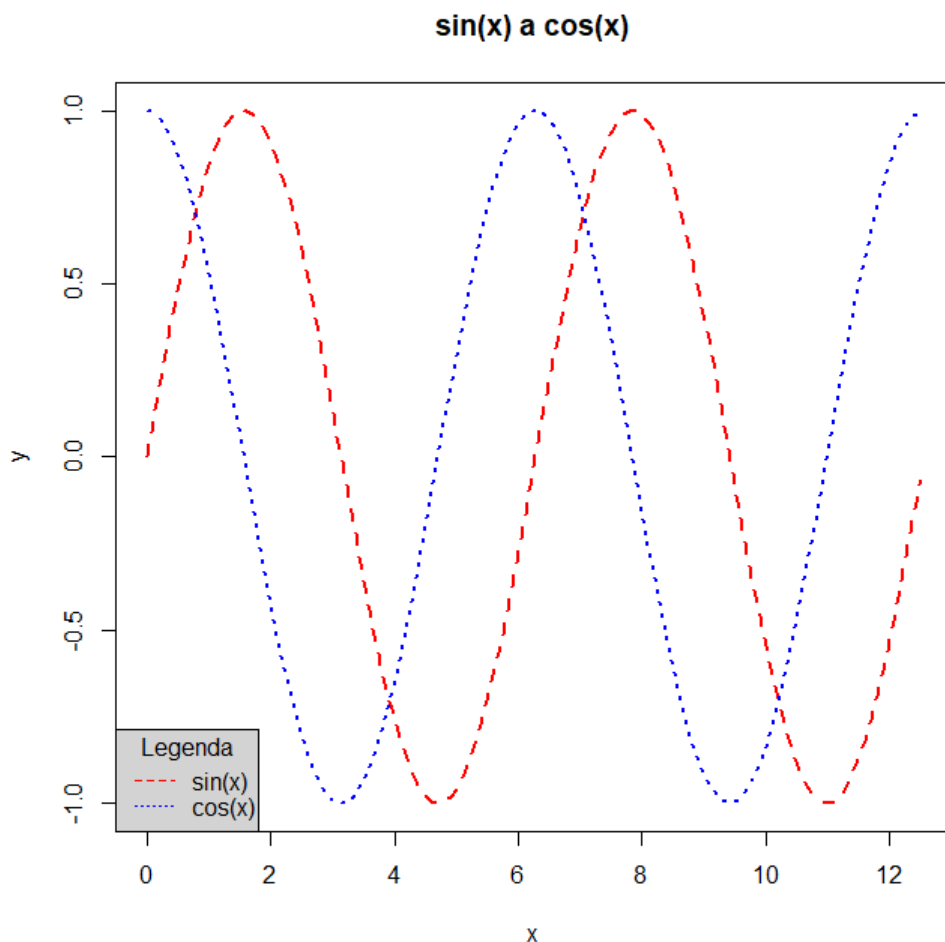
Zdroj: Vlastní zpracování.

Tento způsob editace vnějších okrajů („outer regions“) lze uplatnit jak na grafy vytvořené pomocí parametrů `mfrow/mfcol` z funkce `par()`, tak ale i na `layout()` funkci. Tato metoda nefunguje u grafů vytvořených funkcí `split.screen()`.

6.2.5 Tvorba legendy

Vytváření legendy v jazyku **R** je relativně komplikované. Tato situace souvisí s tím, že legenda není nijak propojená s metodou a atributy pro vizualizaci dat v samotném grafu. Je tedy nutné ručně přiřadit danému popisku správnou barvu a značku. Další problém s legendou souvisí s její pozicí v rámci grafu. Legenda je umístěná uvnitř samotného grafu. Nejprve si vytvoříme standardní legendu pro graf dvou proměnných pomocí funkce `legend()`.

```
1 > x <- seq(from = 0, pi * 4, by = .1)
2 >
3 > # Zakladni graf
4 > plot(
5 +   sin(x) ~ x,
6 +   type = "l",
7 +   col = "red",
8 +   main = "sin(x) a cos(x)",
9 +   ylab = "y",
10 +  xlab = "x",
11 +  lty = 2,
12 +  lwd = 2
13 + )
14 >
15 > # Pridani car
16 > lines(
17 +   cos(x) ~ x,
18 +   type = "l",
19 +   col = "blue",
20 +   lty = 3,
21 +   lwd = 2
22 + )
23 >
24 > # Tvorba legendy
25 > legend(
26 +   "bottomleft",
27 +   title = "Legenda",
28 +   legend = c("sin(x)", "cos(x)"),
29 +   lty = c(2, 3),
30 +   col = c("red", "blue"),
31 +   bg = 'lightgrey'
32 + )
```



Obrázek 6.26: Tvorba legendy grafu

Zdroj: Vlastní zpracování.

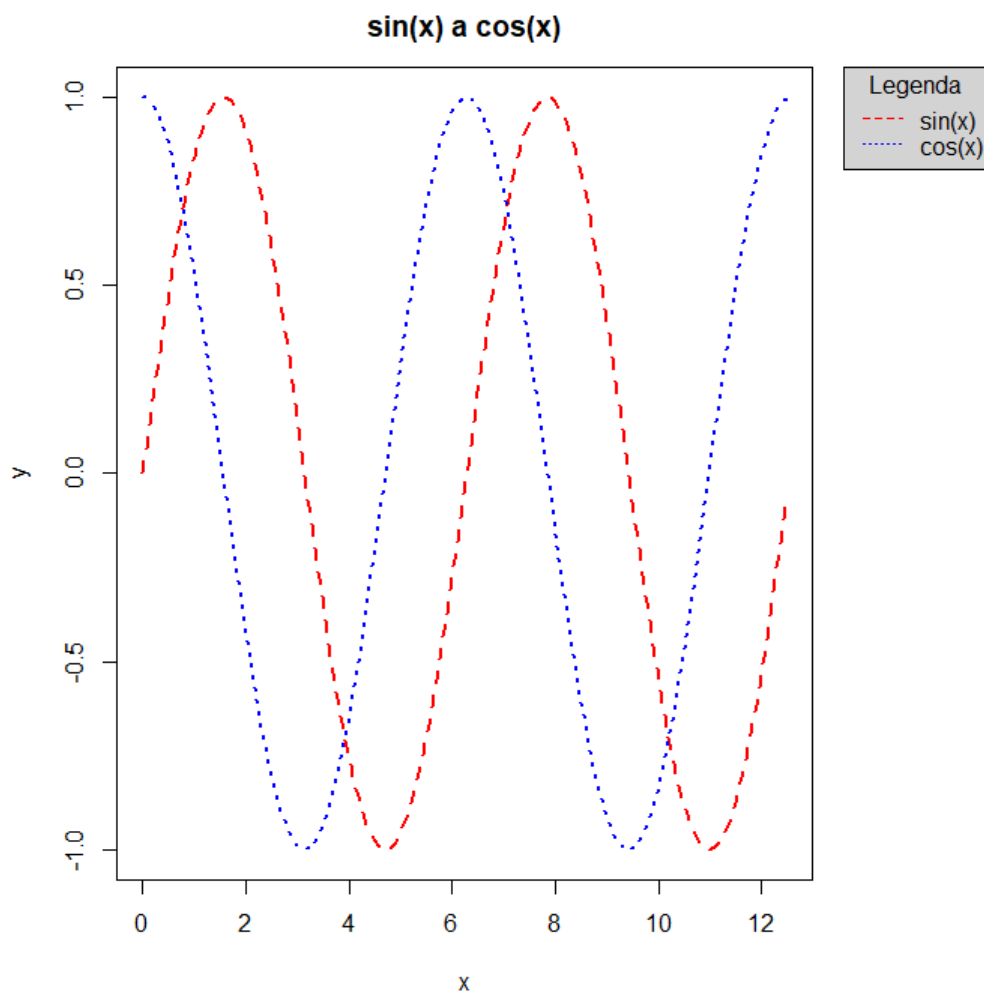
Z příkladu je vidět, že legenda se vytváří stejným způsobem, jako kdybychom vykreslovali graf. Jako první parametr máme pozici legendy. **R** umožňuje textově zadat 9 pozic: „bottomright“, „bottom“, „bottomleft“, „left“, „topleft“, „top“, „topright“, „right“ a „center“. Alternativně lze zadat souřadnice levého horního rohu legendy pomocí parametrů *x* a *y*, které udávají stejné souřadnice, jako jsou na osách grafu pro zobrazení dat. Parametr *title* určuje nadpis legendy. Do parametru *legend* vkládáme jména proměnných (atomický vektor charakterů). V případě, kdy máme spojnicový graf, tak pomocí *lty* vložíme typ čáry. Aby se nám dobře všechny proměnné párovaly, tak je nutné dodržet stejné pořadí (odpovídající popiskům z parametru *legend*). Pokud by data měla být vykreslena pomocí bodů, tak využijeme parametr *pch*. Pomocí parametru *col* zadáme

barvy odpovídající vykresleným proměnným. Parametr `bg` ovlivňuje barvu pozadí grafu. V případě, kdy bychom chtěli vynechat okraje legendy, využijeme parametr `box.lty=0` (přiřadíme šířce boxu 0 bodů). Parametr `box` má i analogie `lwd` a `col` (`box.lwd`, `box.col`).

Jak ale vyřešit situaci, kdy bychom chtěli legendu umístit mimo vnitřní oblast grafu? My sice můžeme legendě dát ručně souřadnice mimo oblast zobrazení dat, ale legenda se mimo tuto oblast nevykreslí. Maximálně bude kus legendy přesahovat dovnitř grafu, ale nic ven. Tato zdánlivě neřešitelná situace má snadné řešení pomocí funkce `layout()` pro více grafů, kdy jsme druhý graf vykreslili bez bodů a os. Toho jsme dosáhli tak, že jsme parametry `xaxt` a `yaxt` nastavili na "n" – žádné okraje. Dále celý okraj (mimo osy i) odstraníme znegováním parametru `frame.plot`.

```
1 > layout(matrix(c(1, 2), ncol = 2), width = c(4, 1))
2 > x <- seq(from = 0, pi * 4, by = .1)
3 > par(mar = c(4, 4, 3, 0))
4 > plot(sin(x) ~ x,
5 +     type = "l",
6 +     col = "red",
7 +     main = "sin(x) a cos(x)",
8 +     ylab = "y",
9 +     xlab = "x",
10 +     lty = 2,
11 +     lwd = 2
12 + )
13 > lines(cos(x) ~ x,
14 +     type = "l",
15 +     col = "blue",
16 +     lty = 3,
17 +     lwd = 2
18 + )
19 >
20 > # Tvorba prazdne grafikce oblasti pro legendu
21 > par(mar = c(4, 0, 3, 1))
22 > plot(c(1),
23 +     type = "n",
24 +     frame.plot = FALSE,
25 +     xaxt = 'n',
26 +     yaxt = 'n',
27 +     xlab = "",
28 +     ylab = ""
29 + )
30 > legend("topright",
31 +     title = "Legenda",
32 +     legend = c("sin(x)", "cos(x)"),
33 +     lty = c(2, 3),
34 +     col = c("red", "blue"),
```

```
35 +     bg = 'lightgrey'  
36 + )
```



Obrázek 6.27: Tvorba legendy okrajů grafu mimo oblast grafu

Zdroj: Vlastní zpracování.

6.2.6 Práce s grafickými okny v softwaru R

V případě, kdy nepracujeme v RStudios ale v základním prostředí jazyka R, se nám nová okna grafů zobrazují v samostatném okně. Zde můžeme snadno nové okno (bez grafu) otevřít pomocí příkazu `dev.new()`, zavřít pomocí `dev.off()`. Pozor – R zavírá aktivní okno! Tedy to, se kterým jsme naposledy pracovali. Příkaz `dev.cur()` vrátí do příkazové

řádky číslo okna, které používáme, a naopak pomocí příkazu `dev.set()` je můžeme přepnout do jiného okna. Příkaz `dev.next()` nám přepne do následujícího grafického okna a příkaz `dev.prev()` do předchozího. Pomocí příkazu `graphics.off()` všechna okna zavřeme. Pokud již máme nějaký graf vytištěný (zobrazený) v samostatném grafickém okně, můžeme toto okno „okopírovat“ pomocí příkazu `dev.copy()` a vložit jej do jiného. Tato situace je výhodná zejména ve chvíli, kdy jsme vytvořili graf a poté jsme si uvědomili, že jej chceme i uložit. Příkaz `dev.print()` vloží okopírované příkazy do nového grafického okna. Výhodou této práce je i to, že můžeme ovlivňovat dopředu to, jak bude grafická oblast velká. Například příkaz `dev.new(height=6, width=3)` vytvoří nové grafické okno o výšce 6 palců a šířce 3 palce.

6.2.7 Jak grafy ukládat

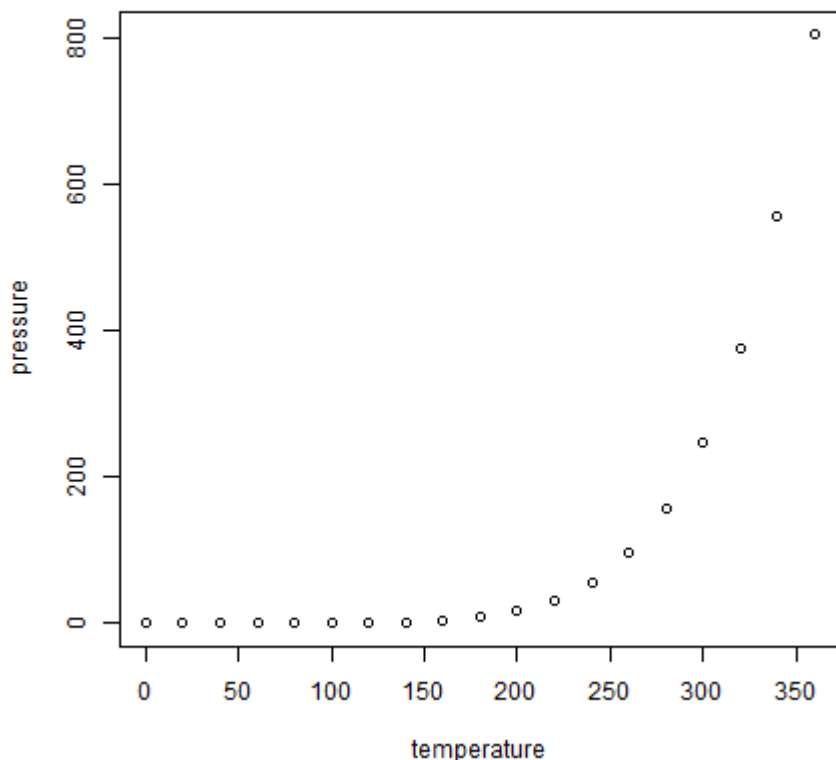
V jazyce **R** máme několik způsobů, jak uložit graf. Úplně nejjednodušší je v RStudiosu kliknout na tlačítko Export a zvolit požadovaný formát, velikost a umístění. V případě, kdy ale potřebujeme ukládání grafů automatizovat, anebo chceme uložit více grafů do jednoho souboru (např. každý graf na jednu stranu souboru pdf), se hodí znát základní příkazy, kterými toho můžeme dosáhnout. Základní způsob ukládání grafu funguje ve třech fázích. V první fázi se vytvoří soubor a otevře se spojení se souborem. Soubor v této fázi vidíme v počítači, jmenuje se tak, jak požadujeme, ale obsahuje 0 bitů dat. Vytvořený soubor je tedy rezervovaný pro naši úlohu. V druhé fázi soubor začínáme naplňovat. Naplňování souboru provádíme tak, jako bychom vytvářeli graf zcela normálně. Voláme tak například funkci `plot()` či `par()` a jiné. V této fázi se všechna data ukládají do dočasné paměti počítače. Graf se nikde nezobrazuje a v **R** se ani neotevře okno. Funkce je tedy nutné si předem připravit a vyzkoušet, protože v této fázi vlastně uživatel nevidí grafický výstup. V třetí fázi pak uzavíráme spojení s grafem (souborem). Při uzavírání spojení dojde k přesunu dat z mezipaměti počítače do daného souboru. Soubor nebude již rezervován softwarem **R** a uživatel si jej může otevřít v jakémkoliv prohlížeči obrázků/fotek/PDF či editoru.

V první fázi je nutné definovat, kde se v počítači soubor bude nacházet, jak se bude jmenovat a jaké bude mít parametry. Za parametry souboru považujeme rozměry obrázku, rozlišení a jiné vlastnosti. První fázi lze provést pomocí funkcí: `bmp()`, `jpeg()`, `png()`, `tiff()`, `pdf()`, `postscript()`, `bitmap()`, `pictex()`, `xfig()`, `win.metafile()`. Jména funkcí napovídají typům souborů, ve kterých chceme graf uložit. Tyto funkce lze rozšířit o další formáty (`devGTK`, `devJava`, `devSVG`...) pomocí dodatečných balíčků. Ukažme si jednoduchý způsob, jak uložit graf vztahu mezi tlakem a teplotou rtuti.

```
1 > getwd()
2 "C:/Users/Documents/R"
3 > #První fáze - otevření spojení se souborem
4 > png(filename = "pressure.png",width=450,height=450)
5 > #druhá fáze - tvorba grafu (popřípadě jakékoli jiné grafické
   funkce)
```



```
6 > plot(pressure)
7 > dev.off() #treti faze - uzavreni spojeni
8 > dir()
9 [1] "pressure.png"
```



Obrázek 6.28: Ukládání grafů

Zdroj: Vlastní zpracování.

Alternativou k tomuto postupu je uložit již zobrazený graf. Pro tento účel má jazyk **R** funkci `savePlot()`. Hlavní parametr této funkce je adresa a jméno souboru. Dále `type` – specifikace typu souboru (`wmf`, `emf`, `png`, `jpg`, `jpeg`, `bmp`, `tif`, `tiff`, `ps`, `eps`, `pdf`). Parametr `device` slouží k identifikaci okna grafu, které se má uložit. Tento postup ale bude fungovat pouze v případě, kdy je graf vytvořen do samostatného okna. Proto v případě RStudio musí uživatel zavolat před tvorbou grafu příkaz `dev.new()` – který budoucí graf vytvoří do nového okna.

```
1 > dev.new()
2 > plot(pressure)
3 > savePlot("pressure2.png")
```

Závěrem k ukládání grafů je dobré například poznamenat, že absolutní/relativní nastavení parametrů ovlivňující rozměry objektů v grafu má velký vliv na výslednou podobu grafu při ukládání. V případě, kdy bude čtenář vyrábět větší množství grafů, je vhodné si předem rozmyslet rozlišení, ve kterém se budou výsledné grafy ukládat, a následně je držet, neboť stejně nastavená tloušťka čáry při rozlišení grafu 450×450 px bude mít úplně jinou podobu v grafu při 900×900 px.

6.3 Ggplot2

Knihovna `ggplot2` je jednou z nejpoužívanějších knihoven pro vizualizaci dat v softwaru **R**. Tato knihovna je postavená na konceptu „grammar of graphics“ (Wickham, 2009; Wickham a kol. 2020), který říká, že jakýkoliv graf může být vytvořen z několika málo komponent (data set, vizuální reprezentace dat a systém souřadnic). Prvenství této knihovny je dáno tím, že se snaží co nejvíce usnadnit práci uživateli a provádět co nejvíce úloh automaticky za něj.

Práce s knihovnou je pro uživatele dost intuitivní, například v `base graphics` jsme složitě přidávali legendu do grafu. Museli jsme definovat ideální polohu, přiřadit grafické atributy a spojit je se jmény proměnných. Knihovna `ggplot` ji vytvoří automaticky a všechny atributy grafu správně sama spáruje. Knihovna disponuje předem definovanou barevnou paletou, která je kontrastní a pomáhá rozeznávat jednotlivé kategorie vizualizovaných dat. Křivky vykreslené touto knihovnou nepůsobí „kostrbatě/hranatě“, tak jako například z knihovny `base graphics`. Knihovna `ggplot` pracuje ve vrstvách. Postupně na sebe vrstvíme jednotlivé komponenty grafu – od surových dat, přes vizualizaci jednotlivých proměnných, až po vizualizaci statistických veličin.

Knihovna má samozřejmě i nevýhody. S ohledem na to, že každý grafický objekt se vizualizuje ve vlastní vrstvě, tak občas vytvoření grafů z velkého množství dat může i pár sekund trvat a může mít větší velikost z pohledu místa na disku. Grafici ale na oplátku ocení uložení dat ve vrstvách a křivkách pro další zpracování. Knihovna není nijak stará, ale i tak již prošla vývojem (proto `ggplot2` a ne `ggplot`). Proto je dobré pro uživatele, kteří mají zájem se dozvědět více o tvorbě grafů v knihovně `ggplot`, aby se nespolehali na starší publikace a vyhledali si zejména aktuální.

6.3.1 Základní model grafu

Hlavní funkcí, kterou budeme v této kapitole používat, je funkce `ggplot`. Knihovna `ggplot2` disponuje ještě funkcí `qplot()` – „quick plot“. Více o této funkci si povíme až v poslední

podkapitole ke knihovně ggplot2.

Jak již bylo zmíněno v úvodu, každý graf má tři základní prvky:

- data,
- aesthetics,
- geometry.

Nejprve musíme mít data, která chceme vizualizovat. Tato data musí být formátu ve `data.frame`. Pokud tomu tak není, tak se je pokusí funkce `ggplot()` na tento formát převést. Knihovna `ggplot2` předpokládá, že uživatel pracuje s čistými daty. Tedy s objektem typu `data.frame`, kde jsou všechny proměnné vhodně pojmenované a data jsou vyčištěná. Při vizualizaci dat se jména proměnných propíšou do popisků grafů, legend.

Aesthetics definuje, jaké atributy (velikost, barvu...) mají vizualizované grafické objekty, které reprezentují data. Proměnné z dat je nutné na tyto vlastnosti namapovat. Toto provádíme pomocí funkce `aes()`, do níž přiřadíme potřebné napojení na data. Funkce `aes()` pak vstupuje do parametru `mapping()`.

Geometry představují grafickou reprezentaci dat, tedy jakým grafickým objektem jsou data prezentovaná (body, čáry, sloupce...). Grafickou interpretaci přiřadíme do grafu pomocí `geoms` funkcí. Všechny funkce pro přidání vrstvy s grafickou interpretací dat začínají `geom_` (popř. na `stats_` - o nich se zmíníme dále). Použití funkce `ggplot()` je následující:

```
1 > ggplot(data = <data>, mapping = aes(<aesthetics>)) + <geoms_funkce>
```

Alternativně se můžeme setkat i s variantou, kdy aesthetics je deklarované („namapované“) uvnitř geom funkce.

```
1 > ggplot(data = <data>) + <geoms_funkce>(mapping = aes(<aesthetics>))
```

Zejména u grafů, které se skládají z různých grafických vrstev, má toto použití své opodstatnění. Nic nám ale nebrání přenést i přiřazení vlastního zdroje dat (parametr `data`) do funkce „geom_“. V případě jednoduchého grafu je ale přehlednější definovat tyto parametry v hlavní funkci `ggplot`. Aesthetic definujeme pomocí funkce `aes()`. Uvnitř této funkce přiřadíme jednotlivé proměnné z dat do příslušných parametrů. Uživatel může využít jakéhokoli způsobu zmiňované deklarace parametrů, ale může i tyto způsoby kombinovat.

V případě `ggplot()` funkce je vhodné ještě rozlišit další důležité prvky vizualizace grafu. Prvním jsou souřadnice (`coordinates`). Není tím myšleno souřadnice ve smyslu geografickém, ale souřadnice bodů v grafu – tedy obvykle kartézské. Funkce, které začínají na

`coord_`, ovlivňují způsob, jakým jsou souřadnice grafu vizualizované. Můžeme použít lineární souřadnice, nelineární souřadnice, fixovat poměr os, zvětšit vybranou část grafu či prohodit osy.

Funkce začínající na `scale_` nám ovlivňují způsob, jakým jsou namapované proměnné vizualizované z pohledu atributů geometrických objektů (barva, velikost...). Pomocí nich můžeme ovlivnit rozsah os, rozdělení škál, ale i přidat jména na jednotlivé osy.

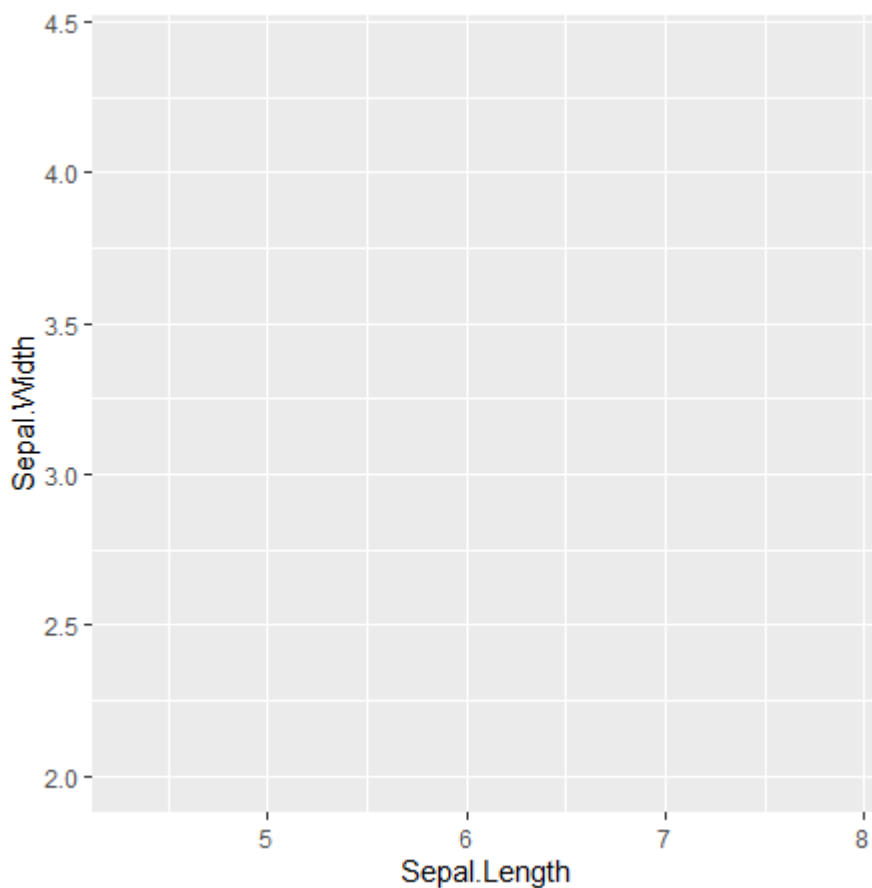
V případě, kdy chceme do grafu přidat grafickou vizualizaci statistické aplikace, využijeme tzv. `stats` funkcí. Jedná se tedy o funkce, které začínají na `stat_`, a jako funkce `geom_` vytváří i funkce `stats` vlastní grafickou vrstvu. I `stats` funkce vytváří vlastní geom objekt, který je vztažený k `aesthetics`. U některých funkcí jsme schopni dosáhnout velmi podobného či až stejného výsledku (`stats` vs `geoms`), avšak primární rozdíl spočívá v tom, že `stats` funkce nejprve statisticky interpretují `aesthetics` a až následně vytvářejí geom objekt.

`Facets` funkce plní velmi podobnou roli jako `mfrow/mfcol` parametr v `par()` funkci u základních grafů (base graphics). Umožní nám tedy rozdělit grafickou oblast na „buňky“, které následně vyplníme jednotlivými grafy.

`Guides` a `themes` funkce nám umožní vytvořit popisky os a upravit další atributy grafu, jako jsou typ a velikost písma. Jak `guides`, tak `theme` funkce vytváří legendu grafu.

Zkusme si udělat základní graf. Ve funkci `ggplot` přiřadíme data a „namapujeme“ jednotlivé proměnné na osy grafu.

```
1 > ggplot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal.  
  Width))
```

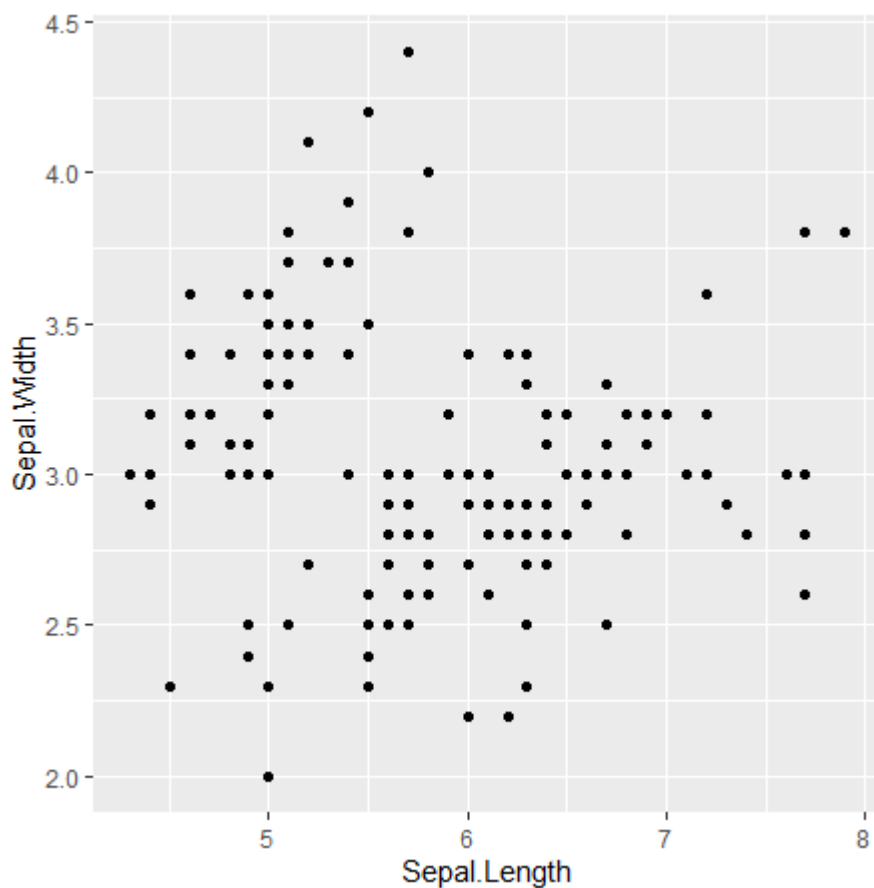


Obrázek 6.29: Funkce ggplot()

Zdroj: Vlastní zpracování.

Pokud postrádáte data v grafu, tak je postrádáte správně – nejsou zde. A to proto, že není přiřazená geom či stat funkce, která by daná data reprezentovala. Zajímavé na tomto grafu je ale to, že jsou zde již přiřazená jména proměnných na obou osách a že funkce ggplot() vyřešila rozsah obou os (limity os).

```
1 > ggplot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal.  
  Width)) + geom_point()
```



Obrázek 6.30: Funkce ggplot: geom_point

Zdroj: Vlastní zpracování.

Grafická vrstva geom_point zobrazí příslušná data prostřednictvím bodů. Zcela identický graf můžeme získat i pomocí stats funkce pro identifikaci unikátních dat, tedy pomocí stats_unique().

```
1 > ggplot(data=iris) + geom_point(mapping = aes(x = Sepal.Length, y =  
Sepal.Width))
```

A jak jsme zmínili na začátku kapitoly, tak namapování aesthetics lze provést i v grafické části (a vlastně tedy i v stat_funkci()). Následující dva příklady vygenerují zcela identický graf, jaký jsme již vykreslili.

```
1 > ggplot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal.  
Width)) + stat_unique()  
2 > ggplot(data = iris) + stat_unique(mapping = aes(x = Sepal.Length,  
y = Sepal.Width))
```

Knihovna ggplot umožňuje práci s grafy jako s objekty. Je tedy možné si uložit graf do proměnné a pak si jej teprve zobrazit. Dále je možné s touto proměnnou pracovat, přičítat k této proměnné další grafické vrstvy či ji jinak upravovat. Předchozí graf můžeme získat i následujícím způsobem:

```
1 > p <- ggplot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal
    .Width))
2 > p + geom_point()
```

Totéž dostaneme i postupným nasčítáním a následně zobrazením.

```
1 > p <- ggplot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal
    .Width))
2 > p <- p + geom_point()
3 > p
```

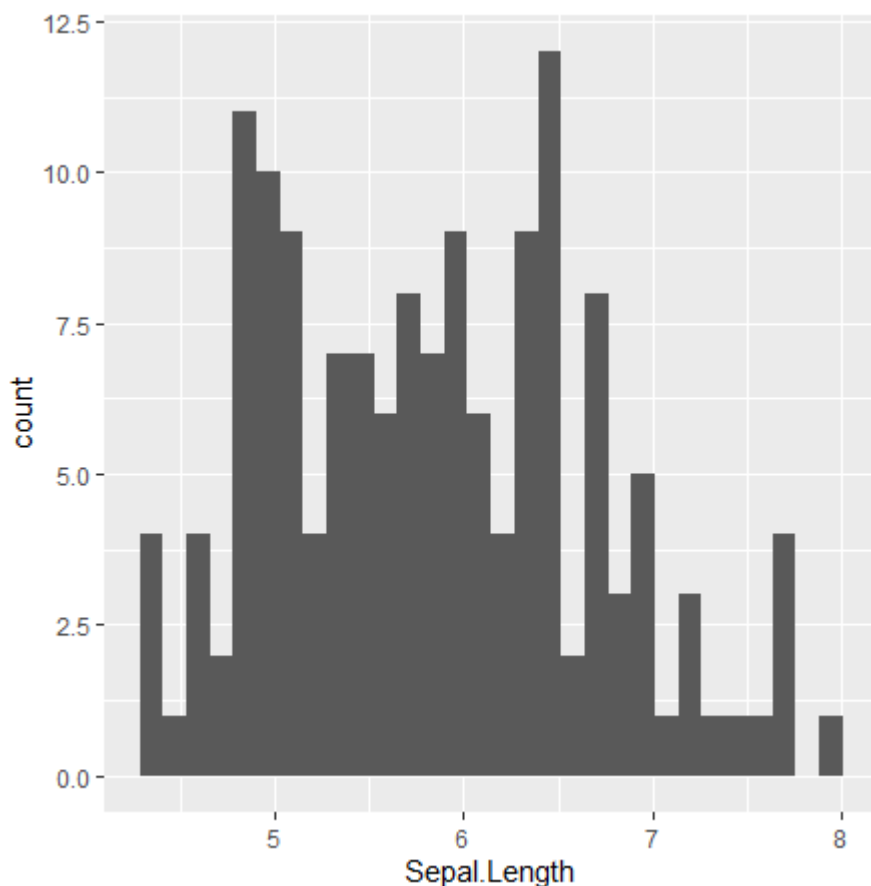
6.3.2 Grafická reprezentace dat – geom a stats

Geoms a stats funkce jsou odpovědné za grafickou vizualizaci dat z pohledu objektu, který je reprezentuje. Definují to, zda objekt bude vizualizován jako čáry, body, sloupce anebo jinou formou. Stat funkce neboli statistická transformace se od geoms odlišuje tím, že před grafickou vizualizací data procházejí transformací, a to obvykle nějakým předem definovaným statistickým postupem. Existuje široká škála geom a stats funkcí. Zaujatému uživateli doporučujeme si projít jednotlivé funkce postupně dle doporučené literatury či dokumentace.

Geoms

V případě grafické reprezentace dat pomocí geoms se obvykle dělí funkce podle počtu proměnných, které zpracovávají. Pro grafickou reprezentaci jedné proměnné máme na výběr funkce, jako je například `geom_area(stat="bin")`. Tato funkce vykreslí plochu grafu pod body, které jsou předané v aesthetics. Funkce `geom_density()` zase vypočte a vykreslí křivku hustoty. Funkce `geom_dotplot()` vykreslí četnostní sloupcový bodový graf dat, který je rozdělený do skupin. Formou obyčejných sloupců (a ne bodů ve sloupcích) můžeme data vykreslit i pomocí funkce `geom_bar()`. Funkce `geom_freqpoly()` a `geom_histogram()` vykreslí rozdělení dat (jejich četnost). První pomocí polygonu (linie) a druhou funkcí pomocí klasického histogramu. Zkusme si teď tedy vytvořit graf histogramu proměnné `Sepal.Length` z datasetu `iris`.

```
1 > p <- ggplot(data = iris, mapping = aes(Sepal.Length))
2 > p <- p + geom_histogram()
3 > p
```



Obrázek 6.31: Funkce ggplot: geom_histogram

Zdroj: Vlastní zpracování.

V případě, kdy vykreslujeme jednorozměrný objekt (jednu proměnnou), není třeba v aesthetic definovat x a y, ale stačí sem vložit tuto proměnnou. Všechny tyto funkce mají mnoho vlastních parametrů, jako jsou například barva (color) či šířky sloupců aj.

Pro dvě spojitě proměnné máme na výběr celou škálu reprezentací dat pomocí geoms. Standardní bodový graf získáme pomocí funkce `geom_point()`. Funkce `geom_jitter()` vykreslí bodový graf, který je „zašumělý“. Body jsou schválně lehce rozhozené od přesného místa jejich polohy. Cílem je zpřehlednit graf v případě, kdy se mnoho bodů zcela překrývá. Kvantily můžeme vykreslit pomocí funkce `geom_quantile()` a pomocí funkce `geom_smooth()` můžeme data protnout vybraným statistickým modelem - například regresí. V následujícím příkladu vykreslíme data `iris` a přidáme vrstvu `geom_smooth()`, která nám umožní přidat regresní funkci do dat.

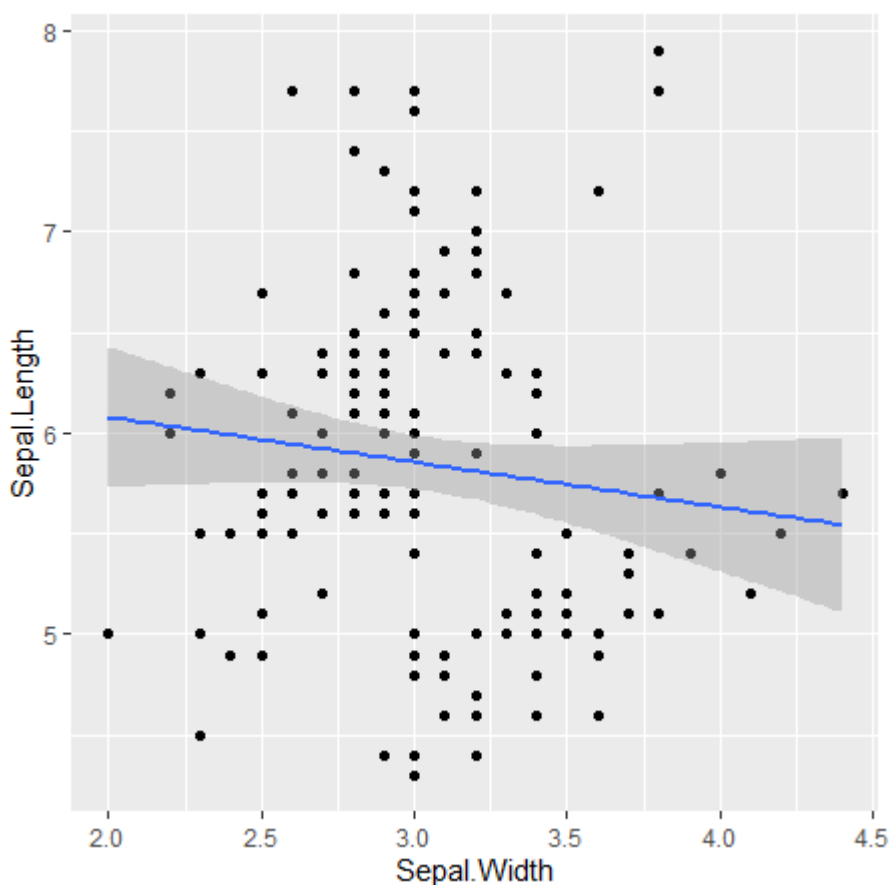
```
1 > p <- ggplot(data = iris, mapping = aes(y = Sepal.Length, x = Sepal
```



```

        .Width))
2 > p + geom_point() + geom_smooth(method = "lm")

```



Obrázek 6.32: Funkce ggplot: `geom_point`, `geom_smooth`

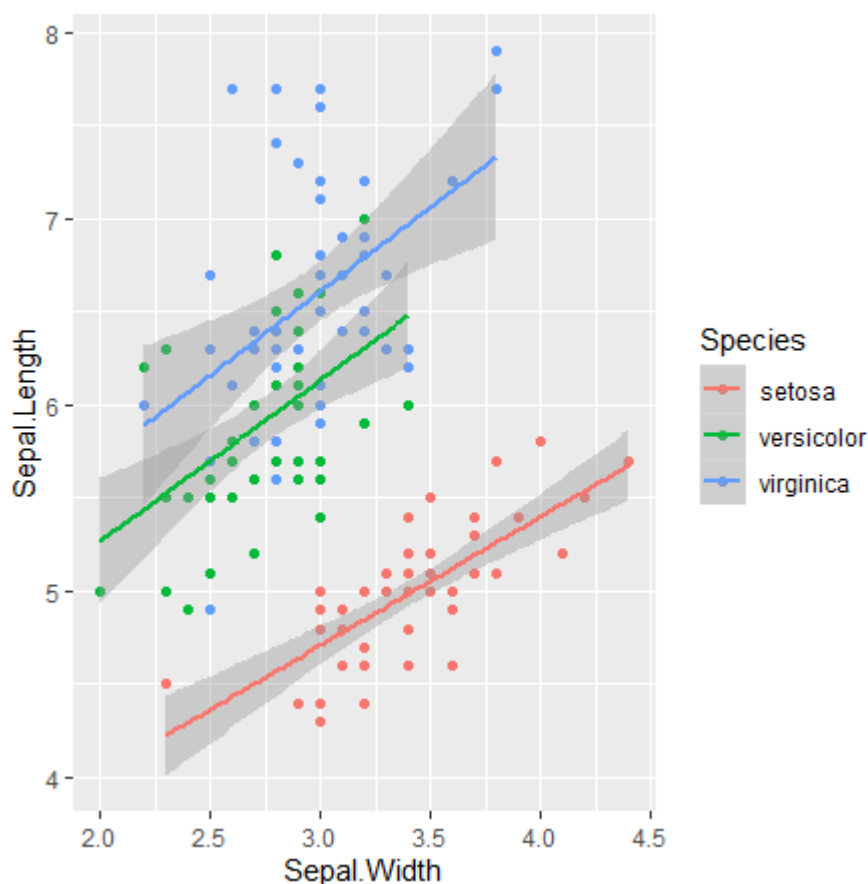
Zdroj: Vlastní zpracování.

Pozorného čtenáře znalého statistických metod jistě daná regrese nepřesvědčí o své vhodnosti. Co kdybychom ale data rozdělili do skupin a následně za ně provedli vykreslení regresních funkcí? V knihovně ggplot to lze snadno. Data můžeme rozdělit do skupin pomocí několika způsobů, a to buď v aesthetics přiřadíme skupiny do proměnné `group`, `fill`, anebo do proměnné `color`. V případě, že přiřadíme skupiny (vektor skupin, vektor faktorů) do parametru `colour`, data budou i rozdílně vykreslena (pomocí barev).

```

1 > p <- ggplot(data = iris,
2 >           mapping = aes(y = Sepal.Length, x = Sepal.Width, color =
3 >           Species))
4 > p + geom_point() + geom_smooth(method = "lm")

```



Obrázek 6.33: Funkce ggplot: geom_point, geom_smooth, pro data rozdělená do skupin

Zdroj: Vlastní zpracování.

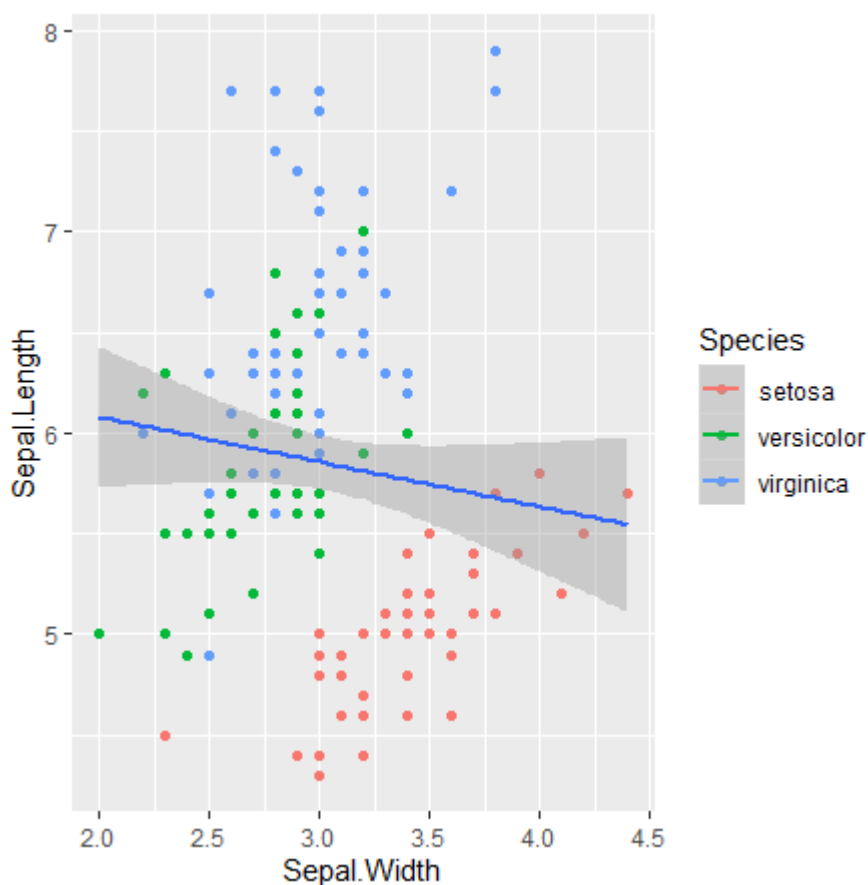
Regresní funkce získané pomocí `geom_smooth()` za jednotlivé skupiny mají hned i jiný sklon. Toto je velmi hezký příklad toho, jak nám může grafická analýza pomoci s pochopením vztahů v datech. Co ale dělat, když chceme mít data rozdělena do barevných skupin, ale výsledek chceme přesto protnout jednou regresní funkcí? Pak nám stačí ve vrstvě `smooth` změnit hodnotu deklarace parametru `skupin`, popř. vyrušit deklaraci barev pro danou vrstvu (pro `geom_smooth()`). Oba následující způsoby vedou ke stejnému výsledku.

```
1 > p <- ggplot(
2 +   data = iris,
3 +   mapping = aes(y = Sepal.Length, x = Sepal.Width, colour =
   Species))
```

```

4 > p + geom_point() + geom_smooth(aes(group = 1), method = "lm")
5 > p + geom_point() + geom_smooth(aes(colour = NULL), method = "lm")

```



Obrázek 6.34: Funkce ggplot: geom_point, geom_smooth pro data celkem

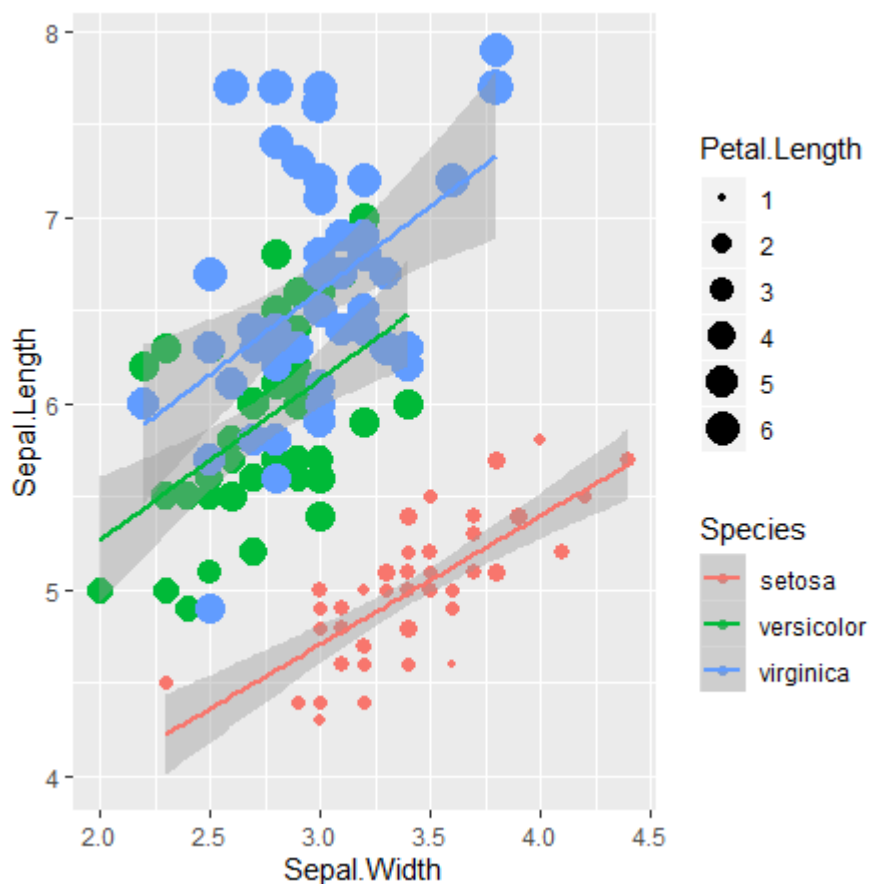
Zdroj: Vlastní zpracování.

Graf s více regresními funkcemi můžeme vylepšit také například tím, že pro jednotlivé body na grafu zvolíme specifické velikosti. Do vrstvy `geom_point()` přiřadíme vlastní aesthetics, kde do atributu velikosti bodů (`size`) vložíme i doposud nepoužitou proměnnou z datasetu `iris`. Volba proměnné je zcela na nás, například `Petal.Length`.

```

1 > p <- ggplot(
2 +   data = iris,
3 +   mapping = aes(y = Sepal.Length, x = Sepal.Width, colour =
4   Species))
5 > p + geom_point(aes(size = Petal.Length)) + geom_smooth(method = "
6   lm")

```

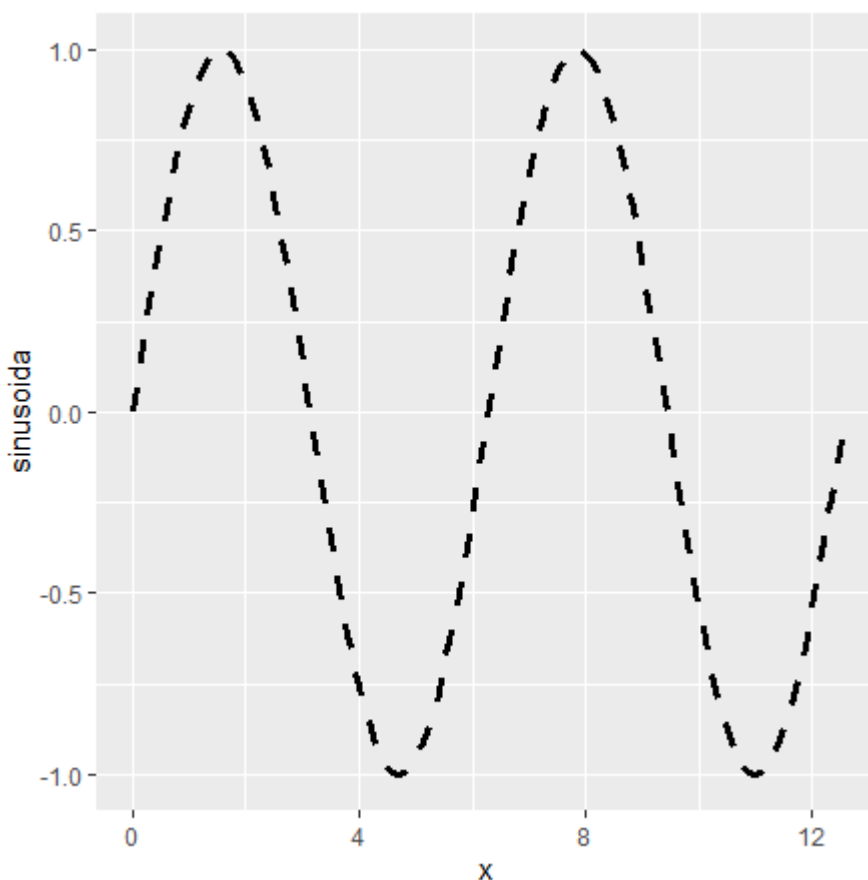


Obrázek 6.35: Funkce ggplot: `geom_point`, `geom_smooth` pro skupiny

Zdroj: Vlastní zpracování.

Za povšimnutí stojí zejména to, že **R** automaticky generuje požadované legendy pro dané třídy dat. Intervaly můžeme zrušit přiřazením hodnoty `FALSE` do parametru `se` ve funkci `geom_smooth()`. Pomocí funkcí `geom_line()` a `geom_area()` můžeme vykreslit liniový graf. V případě funkce `geom_area()` je oblast pod linií vyplněna. Zkusme si jako v base graphics vykreslit i sinusoidu a vykreslit ji do grafu pomocí `geom_line()`.

```
1 > ggplot(data = datasin, mapping = aes(y = sinusoida, x = x)) +
2 +   geom_line(size = 1.2, linetype = 2)
```

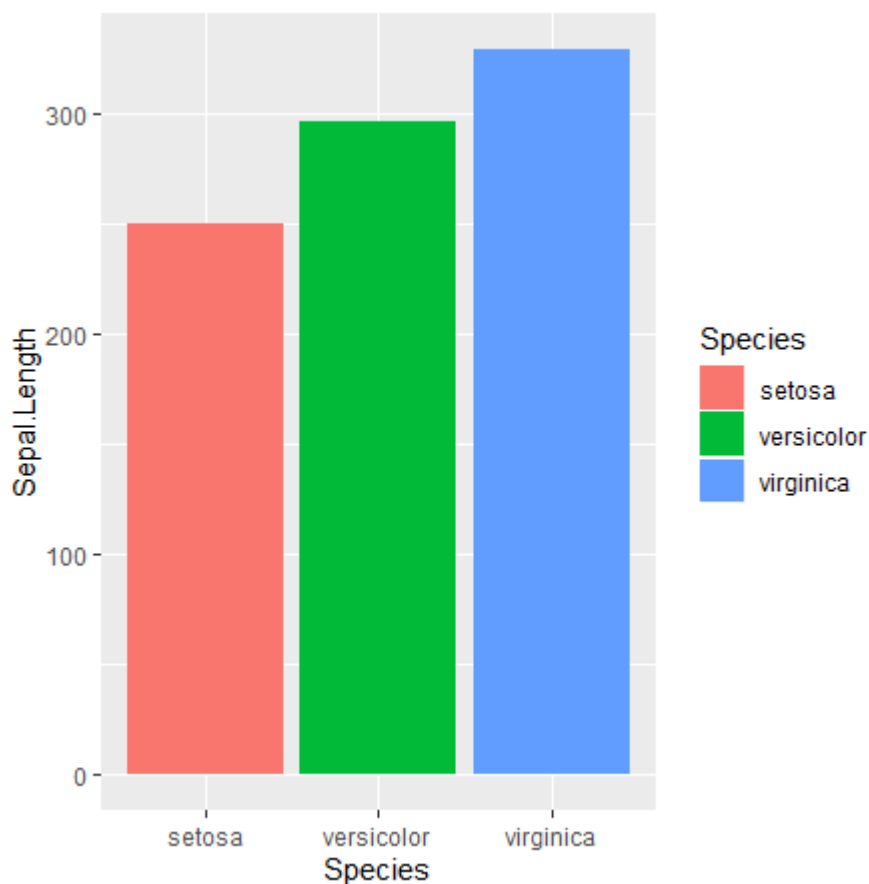


Obrázek 6.36: Funkce ggplot: `geom_line`

Zdroj: Vlastní zpracování.

Pomocí parametrů ve funkci `geom_line()` můžeme nastavit druh, velikost čáry a průhlednost (`alpha`). Klasický sloupcový graf můžeme vytvořit pomocí funkce `geom_bar()`. Data rozdělíme do skupin podle druhu rostliny. Zde je důležité použít místo parametru `group` parametr `fill`. Volba `group` by znamenala rozdělení do daných skupin, ale nevykreslila by sloupce rozdílně. Pokud bychom zvolili argument `colour`, tak by funkce vykreslila okolo každého pozorování rámeček v barvě dané skupiny, ale nikterak by to nepomohlo odlišit barevně dané sloupce. Nastavením parametru `stat` na funkci `identity` získáme klasický sloupcový graf. Funkce `identity` navrací ten samý objekt, který přebírá.

```
1 > p <- ggplot(data = iris, mapping = aes(y = Sepal.Length, x =
  Species, fill = Species))
2 > p + geom_bar(stat = "identity")
```

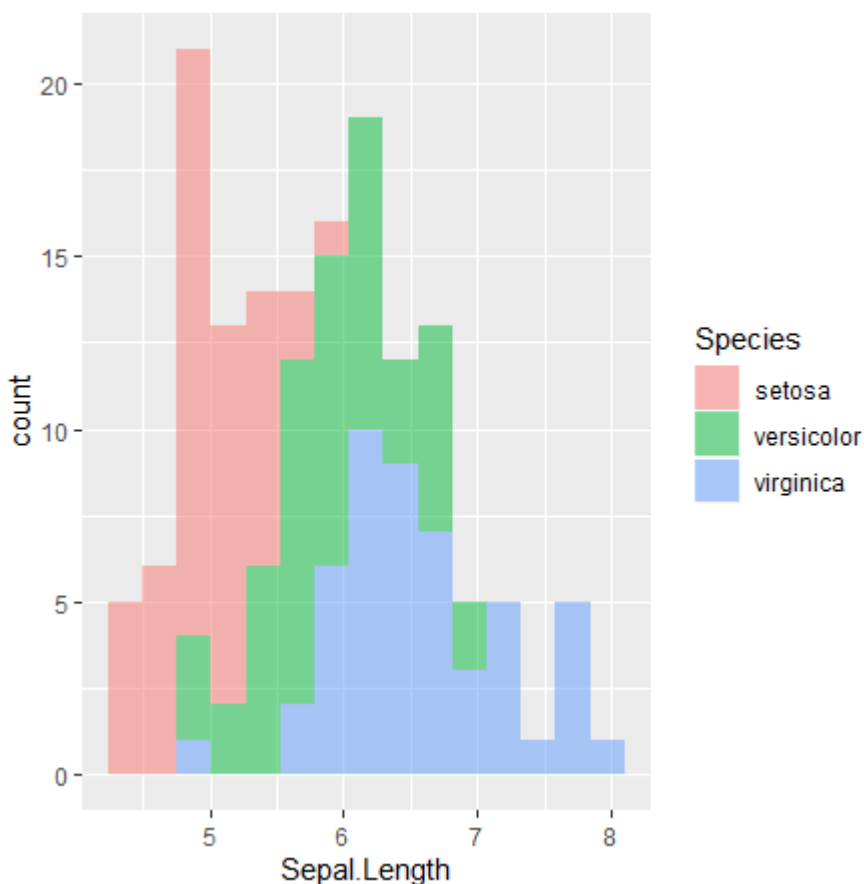


Obrázek 6.37: Funkce ggplot: `geom_bar`

Zdroj: Vlastní zpracování.

Parametr `alfa` definuje míru průhlednosti vizualizovaných dat, například si můžeme zkusit vykreslit tři histogramy pro všechny skupiny `Species` datasetu `iris` do jednoho grafu. Využijeme parametru průhlednosti, abychom zprůhlednili sloupce a bylo vidět „mřížku“ grafu. Dále zvolíme 15 sloupečků (parametr `bins`) pro daný histogram.

```
1 > ggplot(data = iris, mapping = aes(x = Sepal.Length, fill = Species  
2 +   geom_histogram(bins = 15, alpha = 0.5)
```



Obrázek 6.38: Funkce ggplot: geom_histogram

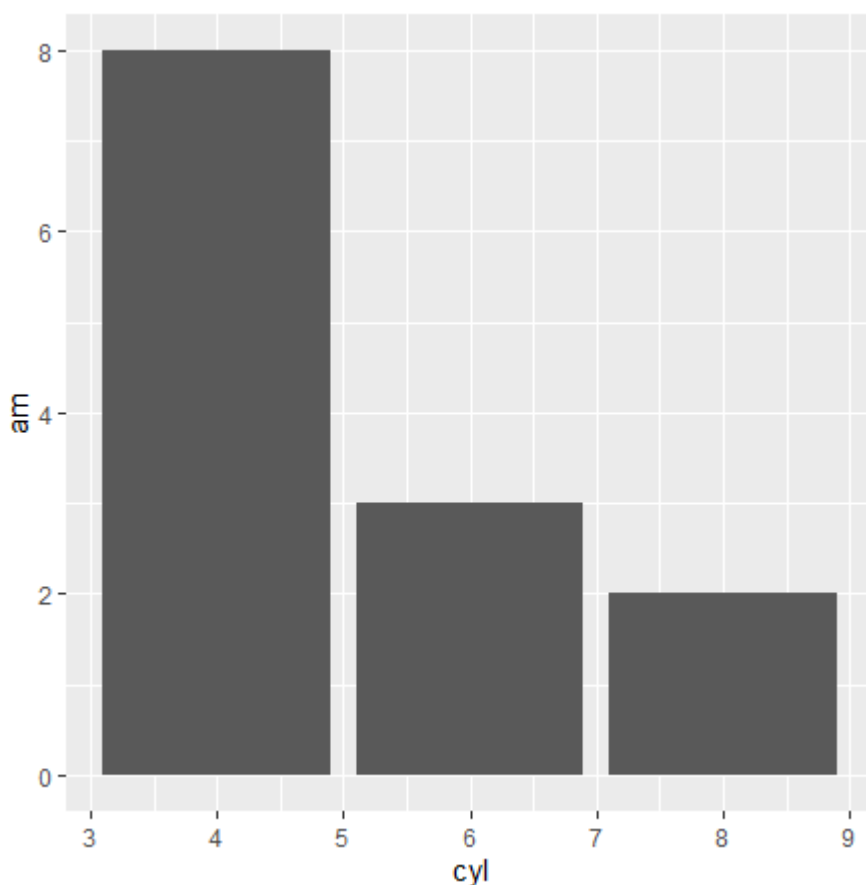
Zdroj: Vlastní zpracování.

6.3.3 Funkce scale

Pomocí `scal_` funkcí můžeme upravovat způsob, jakým jsou aesthetics namapované na jednotlivé grafické vrstvy (geoms). Můžeme ovlivnit barvy, velikosti, ale i rozsah a popisky os (k popiskům se ale dostaneme v následující podkapitole). Po přiřazení proměnných ve funkci `aes()` `ggplot()` sám vybírá nejvhodnější nastavení aesthetics. Způsob výběru těchto nastavení dělá z ggplotu velmi oblíbenou knihovnu – ve výchozím nastavení jsou scales funkce velmi dobře a přehledně nastavené. Avšak může nastat situace, kdy budeme chtít vykreslit některé parametry jinak, popřípadě kdy budeme chtít použít jiné barevné sady či přímo transformovat škálu osy (např. `log`, `sqrt`...). Scale funkce mají obvykle tři společné základní parametry a mnoho dalších, které se vztahují ke konkrétní funkci. Prvním společným parametrem je `name`. Parametr `name` ovlivňuje jméno osy/legendy, ke které se vztahuje daná funkce. Dále se jedná o funkci `limits`. V případě spojitě škály

na dané ose, kterou bychom chtěli transformovat, `limits` přebírá minimální a maximální hodnotu dané osy. V případě kategoriální/diskrétní osy pak tento parametr přebírá vektor kategorií, které má vykreslit. Třetím parametrem je `break`. Tento parametr ovlivňuje zlomy na dané ose či v legendě. Lze jej kombinovat s parametrem `labels`, který definuje, jaké popisky se mají u zlomů vizualizovat. Následující sloupcový graf vizualizuje počty aut s automatickou převodovkou (`am=1`) pro počty válců. Co je na tomto grafu špatně?

```
1 > p <- ggplot(mtcars, aes(y = am, x = cyl))  
2 > p + geom_bar(stat = "identity")
```



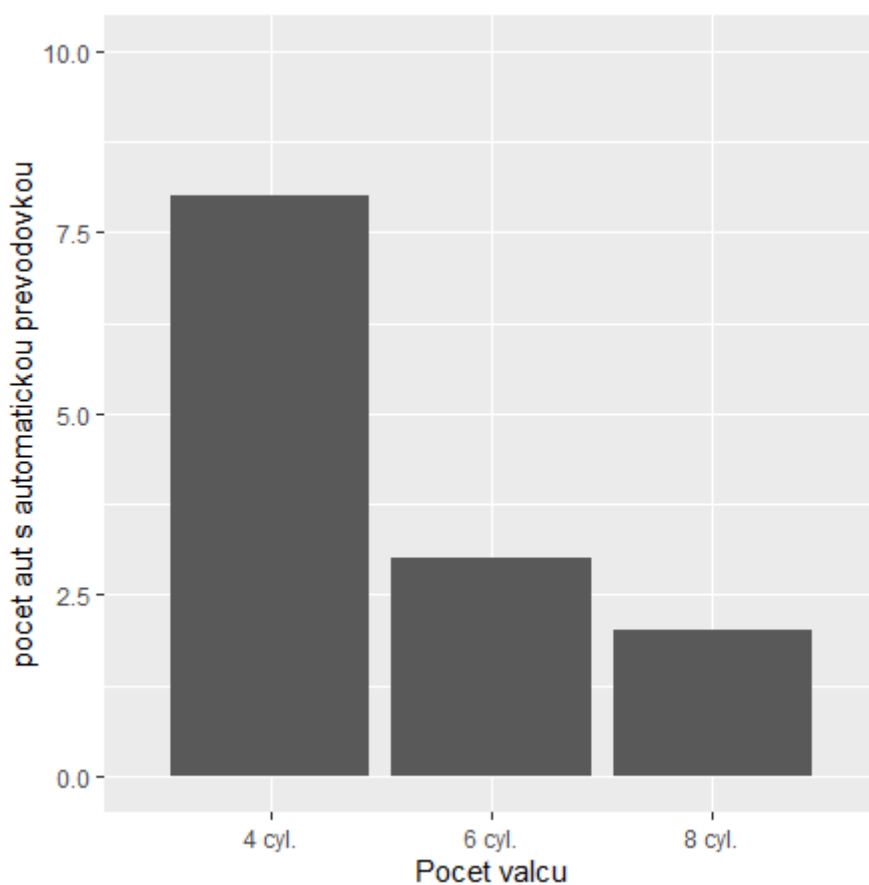
Obrázek 6.39: Funkce ggplot: scale funkce I

Zdroj: Vlastní zpracování.

Možná jste si všimli, že popisky na ose x jsou spojitě. Tato osa je chápána jako spojitá veličina. Jsou ale počty válců na motorech aut spojitě? Nejsou. V tomto případě máme dvě možnosti, jak daný graf upravit. Buď předáme proměnnou `cyl` jako faktor již v aesthetics,

anebo použijeme funkce `scale`. Pro převod dané osy na diskrétní můžeme využít funkce: `scale_x_discrete()` a `scale_y_discrete()`. Naopak pro převod na spojitou můžeme využít funkce: `scale_x_continuous()` a `scale_y_continuous()`.

```
1 > p <- ggplot(mtcars, aes(y = am, x = cyl))
2 > p + geom_bar(stat = "identity") +
3 +   scale_x_discrete(
4 +     name = "Pocet valcu",
5 +     limits = unique(mtcars$cyl),
6 +     labels = paste(unique(mtcars$cyl), "cyl.") +
7 +     scale_y_continuous(name = "pocet aut s automatickou
8 +       prevodovkou",
9 +       limits = c(0, 10))
```

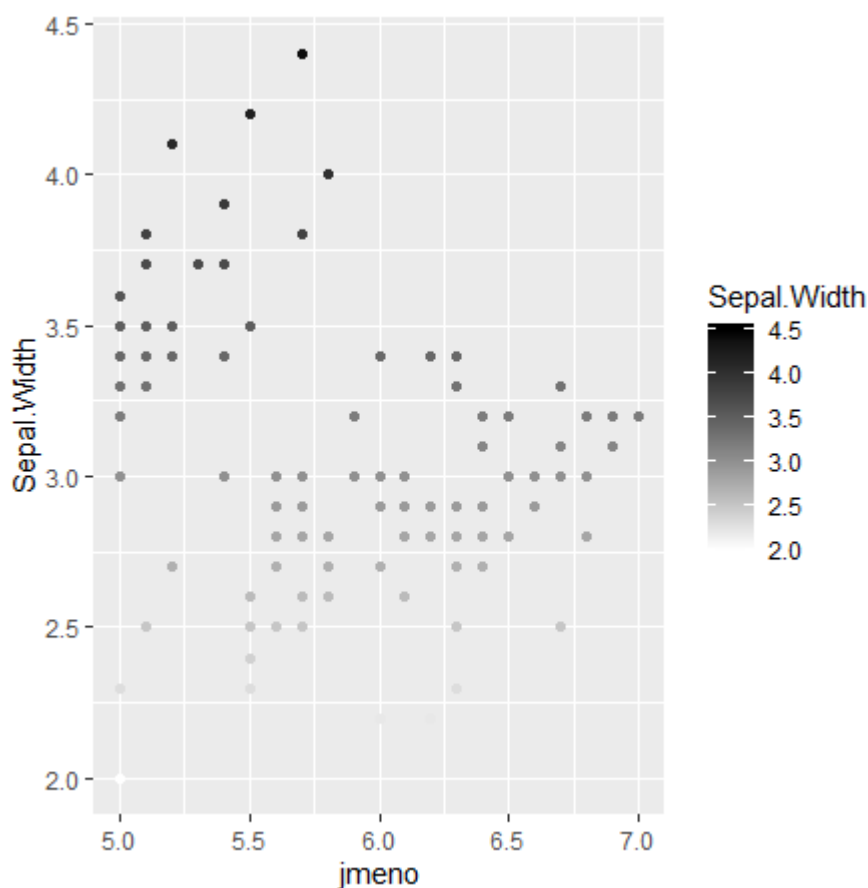


Obrázek 6.40: Funkce ggplot: scale funkce II

Zdroj: Vlastní zpracování.

Obdobně můžeme editovat i barvy, avšak musíme je volit přes parametr, kterým jich bylo dosaženo. Pokud to bylo přes `colour`, tak použijeme funkci `scale_colour()`, pokud to bylo přes `fill`, tak funkci `scale_fill()`. Záleží vždy na specifické situaci a způsobu tvorbu grafu. Barva bodů grafu je předána pomocí parametru `colour` ve funkci `ggplot`. Proto dále volíme funkci `scale_colour()` na změnu barvy. Ve funkci `scale_colour_gradient()` volíme minimální a maximální barvu. Tato funkce existuje ve verzi i pro parametr `fill`.

```
1 > ggplot(data = iris,  
2 +   mapping = aes(x = Sepal.Length, y = Sepal.Width, colour = Sepal.  
   Width)) +  
3 +   geom_point() + scale_colour_gradient(limits = c(2, 4.5),  
4 +                                       low = "white",  
5 +                                       high = "black") +  
6 +   scale_x_continuous("jmeno", limits = c(5, 7))
```

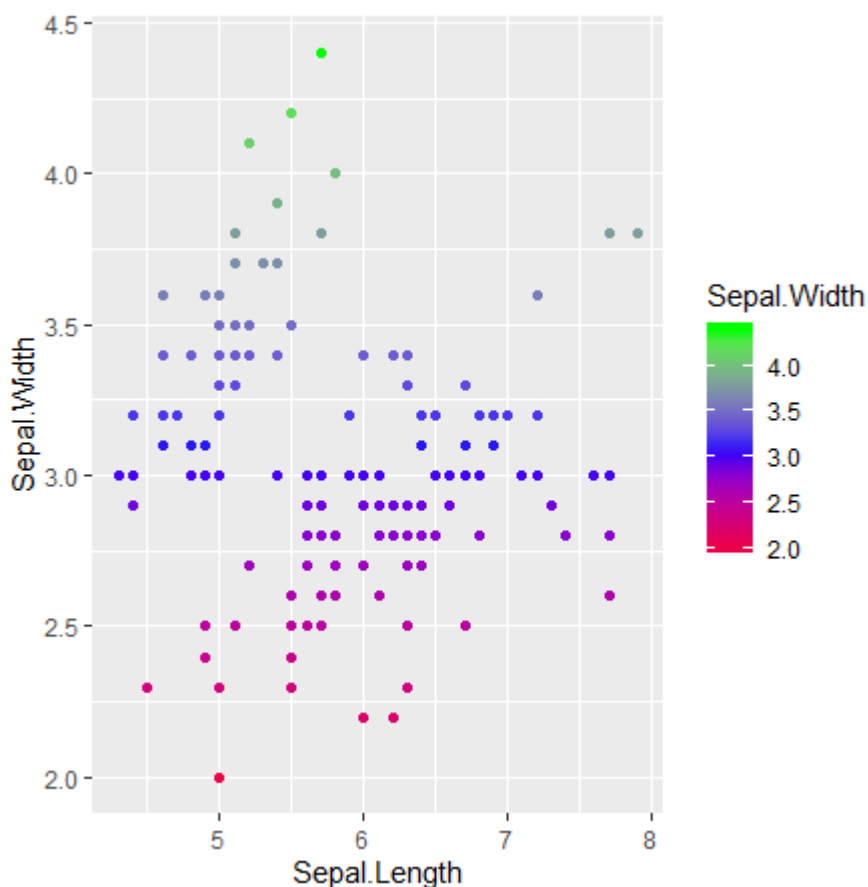


Obrázek 6.41: Funkce ggplot: scale funkce III

Zdroj: Vlastní zpracování.

Ggplot disponuje i funkcemi pro vícepřechodové škály. Tedy pro dva přechody: `scale_colour_gradient` a `scale_fill_gradient2()`. Dále pro n-tý počet přechodů mezi barvami můžeme použít funkce: `scale_colour_gradientn()` a `scale_fill_gradientn()`.

```
1 > plot(data = iris, mapping = aes(x = Sepal.Length, y = Sepal.Width,  
2 +   colour = Sepal.Width))+ geom_point() +  
3 +   scale_colour_gradient2(midpoint = mean(iris$Sepal.Width),  
4 +   low = "red", mid = "blue", high = "green")
```



Obrázek 6.42: Funkce ggplot: scale funkce IV

Zdroj: Vlastní zpracování.

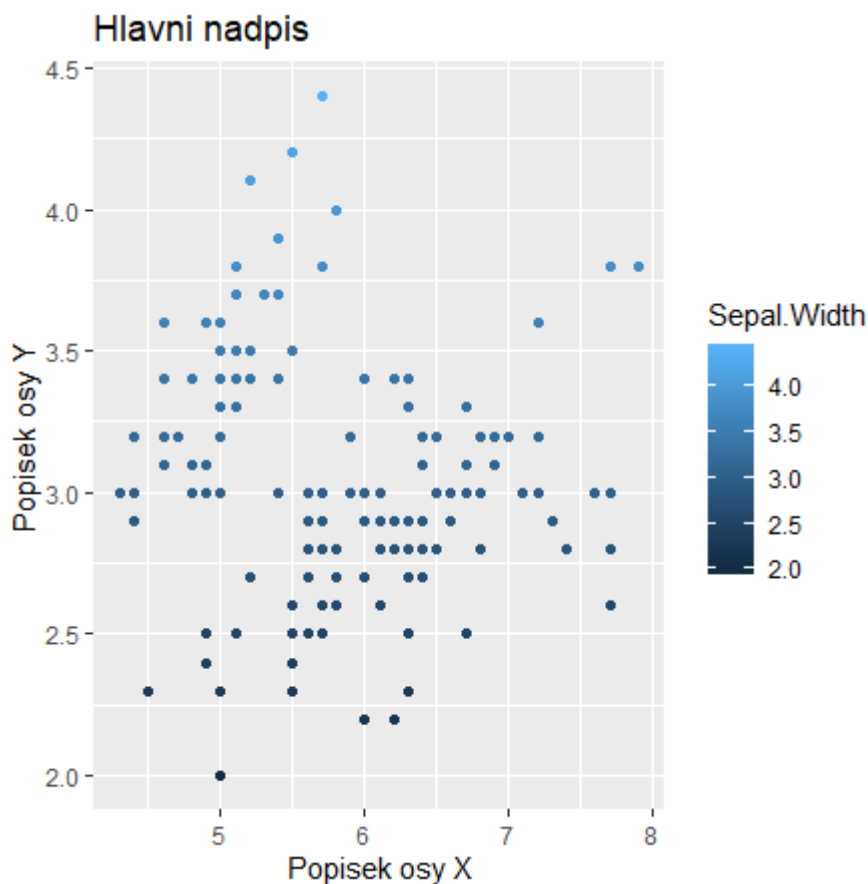
6.3.4 Nadpisy a popisky os

Základní editaci nadpisů a popisků grafu můžeme provést pomocí funkcí: `ggtitle()` (hlavní nadpis), `xlab()` a `ylab()` (popisky os).

```

1 > p <- ggplot(
2 +   data = iris,
3 +   mapping = aes(x = Sepal.Length, y = Sepal.Width, colour = Sepal
4 +   ) + geom_point()
5 > p + ggtitle("Hlavni nadpis") + xlab("Popisek osy X") + ylab("
   Popisek osy Y")

```

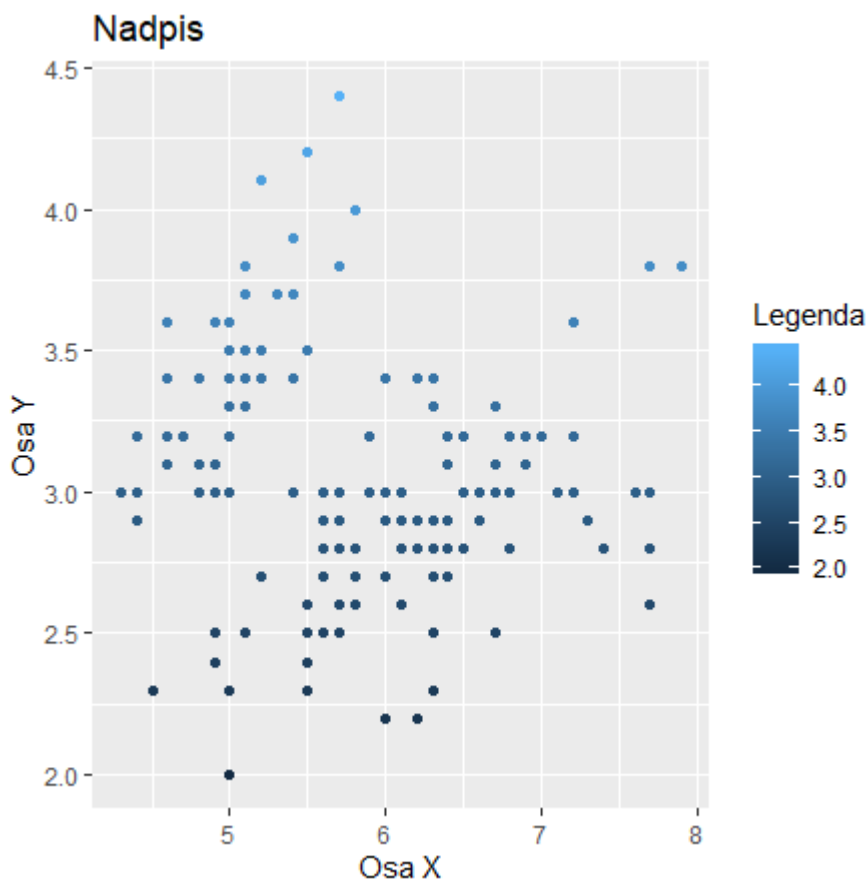


Obrázek 6.43: Funkce ggplot: ggtitle, xlab a ylab

Zdroj: Vlastní zpracování.

Alternativně lze použít funkci `labs()`, která přiřazuje všechny popisky grafu najednou. V případě, kdy chceme editovat legendu či popisky skupin, do kterých jsou data rozdělená, tak to provedeme přes příslušný způsob namapování aesthetics. Pro editaci popisků legendy zde lze použít parametr `legend`.

```
1 > p + labs(title = "Nadpis", x = "Osa X", y = "Osa Y", colour = "
  Legenda")
```



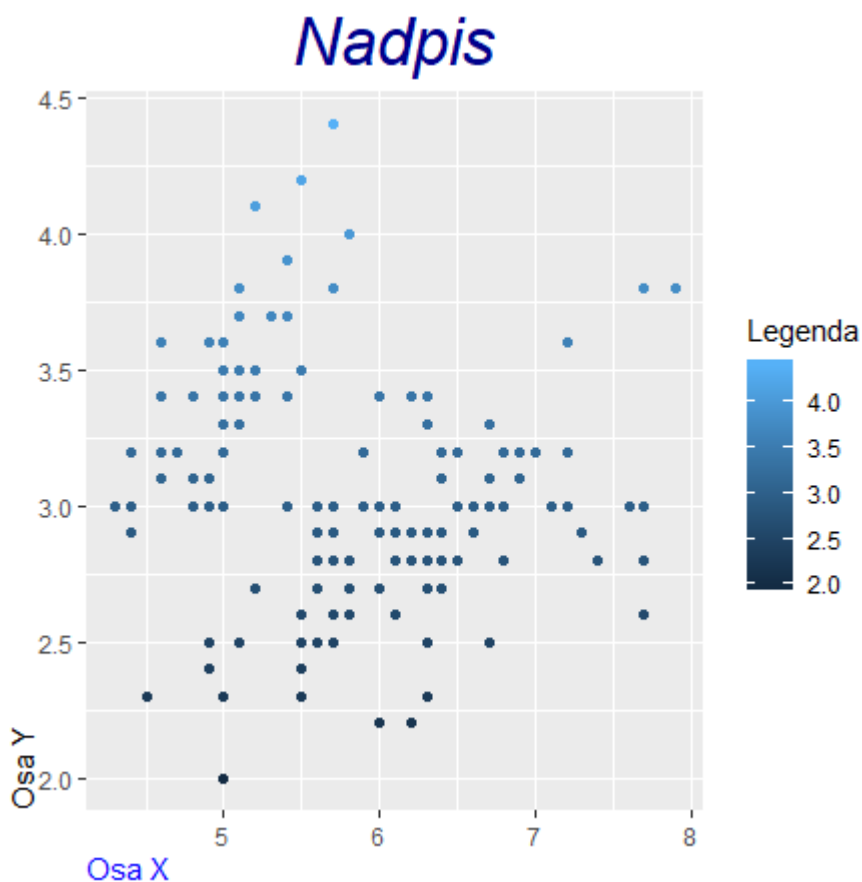
Obrázek 6.44: Funkce ggplot: labs

Zdroj: Vlastní zpracování.

Nejkomplexnější způsob editace atributů grafu lze ale provést pomocí funkce **theme**. Pomocí **theme** můžeme editovat typ, barvu, velikost popisků os. Jejich naklonění, barvu pozadí a ohraničení. Funkce **theme** nám snadno umožní si vytvořit vlastní styl grafu a pak jej můžeme lehce přenášet i na jiné typy grafů. Přes **theme** můžeme editovat velké množství parametrů. Ukažme si základní editaci barev popisků grafu.

```
1 > p <- ggplot(
2 +   data = iris,
3 +   mapping = aes(x = Sepal.Length, y = Sepal.Width, colour = Sepal.
4 +   Width)
5 + ) + geom_point()
```

```
5 > p <- p + labs(  
6 +   title = "Nadpis",  
7 +   x = "Osa X",  
8 +   y = "Osa Y",  
9 +   colour = "Legenda"  
10 + )  
11 > p <- p + theme(plot.title = element_text(  
12 +   colour = "darkblue",  
13 +   size = 25,  
14 +   face = "italic",  
15 +   hjust = 0.5  
16 + ))  
17 > p <- p + theme(axis.title.x = element_text(colour = "blue", hjust  
18 +   = 0))  
18 > p <- p + theme(axis.title.y = element_text(colour = "black", hjust  
19 +   = 0))  
19 > p
```



Obrázek 6.45: Funkce ggplot: theme

Zdroj: Vlastní zpracování.

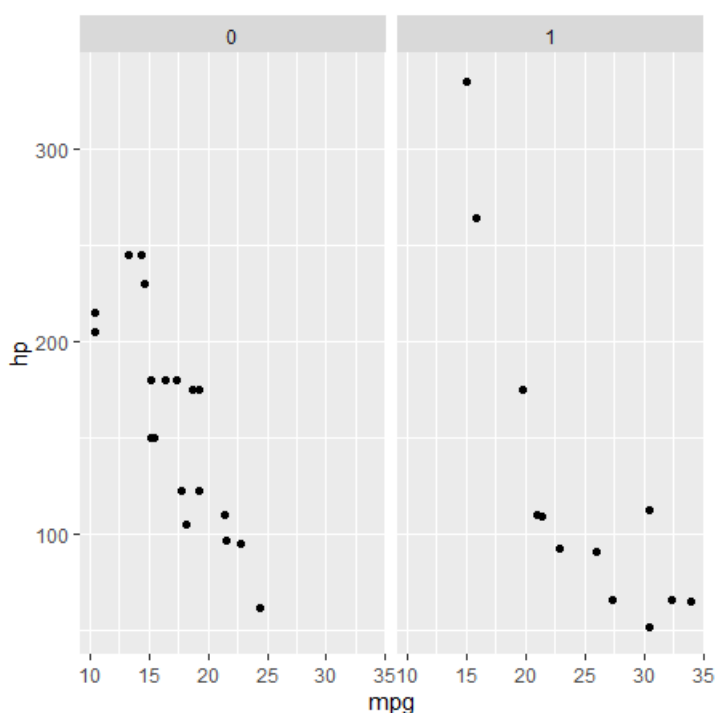
Pomocí parametrů `plot.title` a `axis.title` jsme nastavili formát pro jednotlivé popisky grafu. Pro jejich editaci je nutné využít funkcí `element_`. Každý typ parametru z funkce `theme` má přiřazen druh vstupu. U některých se jedná přímo o specifickou třídu objektu, u jiných je možné více druhů vstupů. V našem případě je pro editaci typu textu nutné použít funkce `element_text`.

6.3.5 Více grafů v jedné grafické oblasti

Knihovna ggplot disponuje podobným způsobem tvorby více grafů v jedné grafické oblasti jako base graphics. Způsob rozdělení grafů je velmi podobný „mřížce“ či „buňkám“, se kterými jsme pracovali pomocí pár funkce v base graphics. Pro tuto úlohu disponuje ggplot dvěma funkcemi. Jedná se o funkci `facet_grid` a `facet_wrap`. První funkce `facet_grid` rozděluje data do 2D pole podle vzorce. Proměnná před tildou definuje počet řádek a

proměnná za tildou definuje počet sloupců, do kterých se grafy rozdělí. Zkusme si v následujícím příkladu vykreslit bodový graf pro data `mtcars` (`mpg` a `hp`) podle druhu převodovky do sloupců.

```
1 > p <- ggplot(data = mtcars, mapping = aes(x = mpg, y = hp))
2 > p + geom_point() + facet_grid(. ~ am)
```

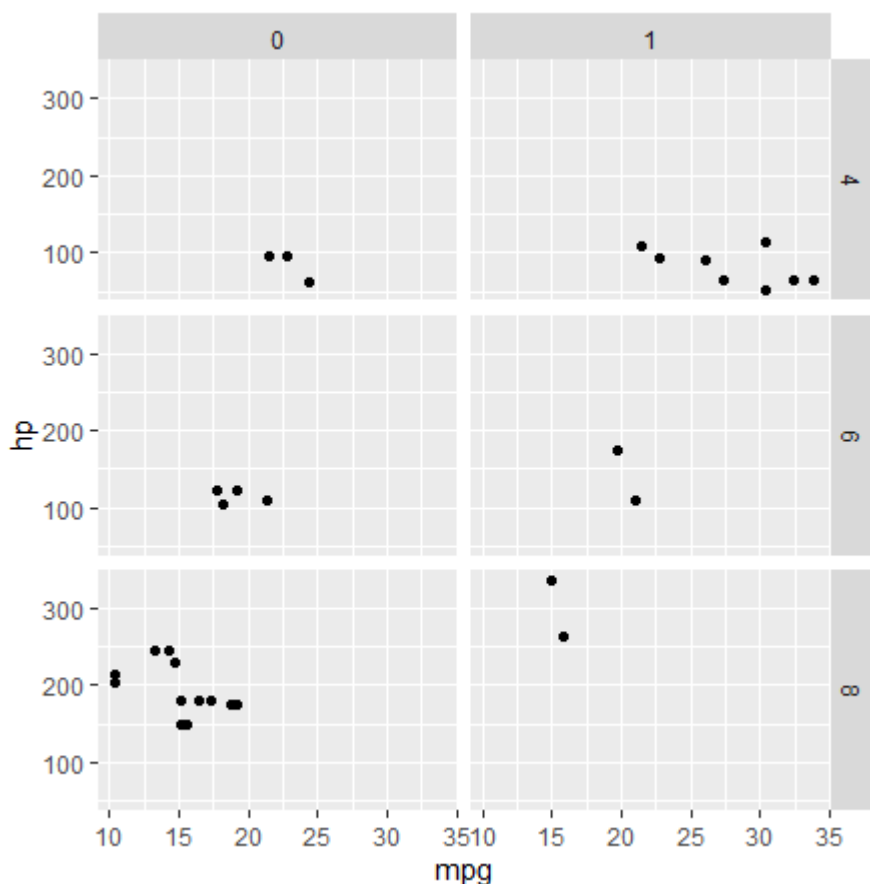


Obrázek 6.46: Funkce `ggplot`: `facet_grid` funkce I

Zdroj: Vlastní zpracování.

Alternativně můžeme data rozdělit jak pomocí proměnné `am` (druh převodovky), tak i pomocí proměnné `cyl` (počet válců motoru auta).

```
1 > p <- ggplot(data = mtcars, mapping = aes(x = mpg, y = hp))
2 > p + geom_point() + facet_grid(cyl ~ am)
```

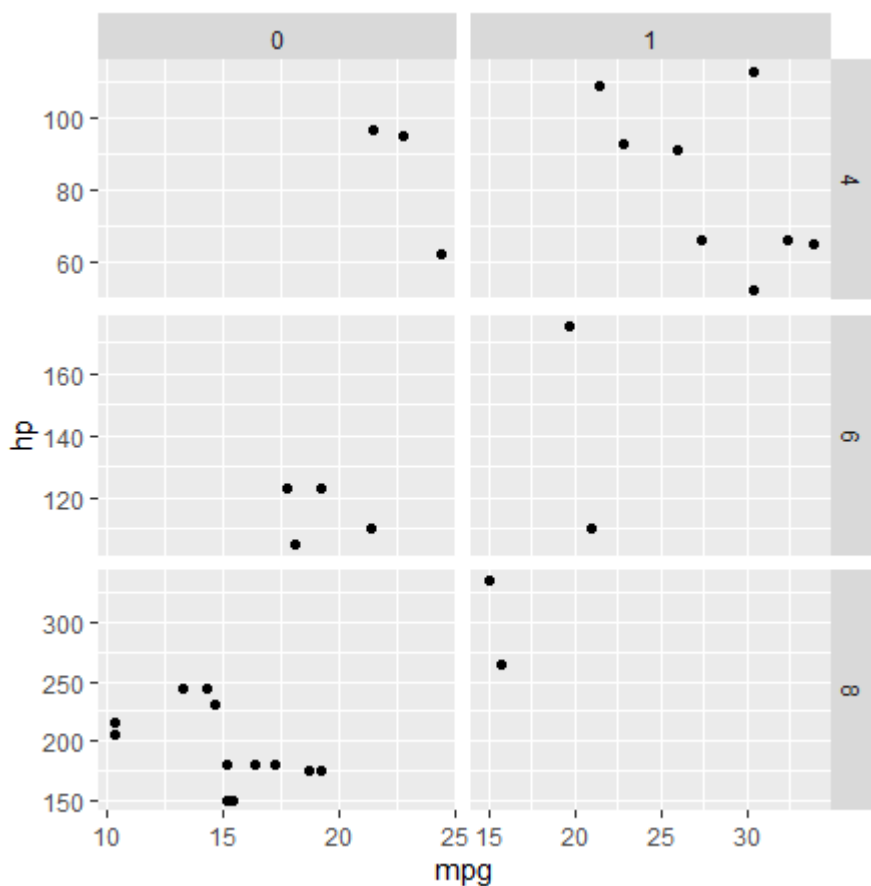



Obrázek 6.47: Funkce ggplot: facet_grid funkce II

Zdroj: Vlastní zpracování.

Jistě jste si všimli, že všechny dílčí grafy mají stejné osy. To se dá ovlivnit pomocí parametru `scales`, který má výchozí parametr „fix“. Na všech grafech jsou stejné rozsahy maxim/minim na osách. Pokud bychom do parametru vložili parametr „free“, tak budou všechny osy rozdílné (svázané jsou jen osy x v sloupcích a osy y v řádkách). Dalšími parametry jsou `free_x` a `free_y`, které uvolňují jen specifickou osu.

```
1 > p <- ggplot(data = mtcars, mapping = aes(x = mpg, y = hp))
2 > p + geom_point() + facet_grid(cyl ~ am, scales = "free")
```

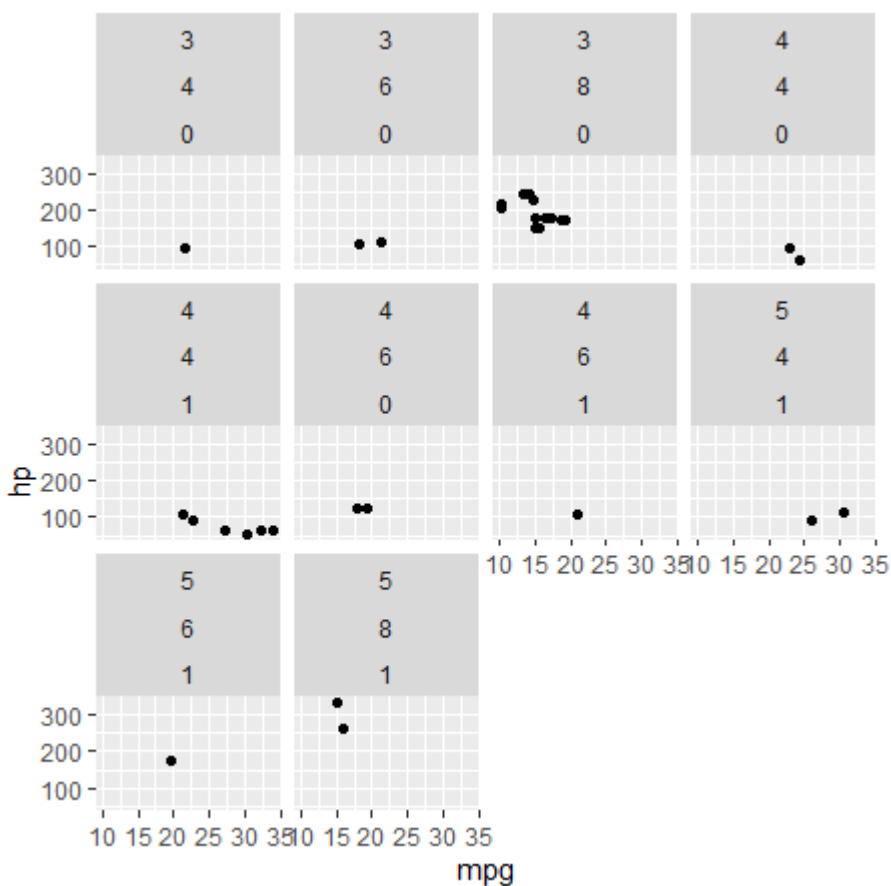


Obrázek 6.48: Funkce ggplot: facet_grid funkce, scales=free

Zdroj: Vlastní zpracování.

Funkce `facet_wrap()` negeneruje cíleně grafy do mřížky podle daných parametrů, ale do panelu grafů. Pro tento panel můžeme zadat i maximální počet sloupců/řádků (`ncol` a `nrow`). Výhodou funkce `facet_wrap()` je to, že můžeme data rozdělit i podle více než dvou proměnných.

```
1 > p <- ggplot(data = mtcars, mapping = aes(x = mpg, y = hp))
2 > p + geom_point() + facet_wrap(gear ~ cyl + am)
```



Obrázek 6.49: Funkce ggplot: facet_wrap funkce

Zdroj: Vlastní zpracování.

6.3.6 Ukládání grafu

Ukládání grafů je v ggplot knihovně velmi podobné jako v případě base graphics grafů. Oproti této knihovně je ale nutné tzv. provést „tisk“ daného grafu pomocí funkce `print()`. Ukážeme si tento postup na grafu datasetu `iris`.

```
1 > p <- ggplot(
2 +   data = iris,
3 +   mapping = aes(x = Sepal.Length, y = Sepal.Width, colour = Sepal.
4 +   ) +
5 +   geom_point()
6 >
7 > png(filename = "GGplot.png",
```

```
8 +     height = 450,  
9 +     width = 450)  
10 > print(p)  
11 > dev.off()
```

Alternativou k tomuto postupu je již zobrazený graf přímo uložit pomocí funkce `ggsave()`. Hlavní parametr funkce `ggsave()` je jméno grafu. Funkce si ze jména grafu přebere sama druh souboru, do kterého ho má uložit, a vybere tak správnou proceduru.

```
1 > ggplot(data = iris,  
2 +   mapping = aes(x = Sepal.Length, y = Sepal.Width, colour = Sepal.  
   Width)) + geom_point()  
3  
4 > ggsave("ggplot.pdf")  
5 Saving 5.53 x 7.77 in image
```

Tato funkce ale funguje jen pro zobrazené grafy. Je tedy nutné si v RStudios/R environmentu graf nejprve zobrazit do nového grafického okna. Ve funkci `ggsave()` můžeme specifikovat množství atributů o daném grafu, který ukládáme, například jeho výšku, šířku, poměr stran, rozlišení či změnit okno, ze kterého se má graf uložit, pokud máme otevřeno více grafických oken.

Kapitola 7

Funkce merge a knihovna dplyr

V rámci této kapitoly si představíme funkci `merge()` a knihovnu `dplyr`. Funkce `merge()` slouží ke spojování datových souborů. Knihovna `dplyr` je pak komplexní knihovnou funkcí, které slouží k pokročilé práci s daty (více viz. Spector 2008).

7.1 Funkce merge

Často nastává situace, kdy potřebujeme dva daty spojit podle určitého klíče. Ve výchozí instalaci softwaru **R** je k dispozici funkce `merge()`. Tato funkce má více parametrů, které mají nastavené důležité výchozí parametry. Postupně si zkusíme tuto funkci projít. Nejprve si vytvoříme dva testovací daty ve formátu `data.frame`. První dat obsahuje seznam pěti jmen žáků a jejich počet bodů ze zeměpisu.

```
1 > data.A <- data.frame(  
2 +       jmeno = c("Karel", "Pavel", "Jan", "Aneta", "Jakub"),  
3 +       zemepis = c(20, 10, 12, 43, 20),  
4 +       stringsAsFactors = FALSE)  
5 > print(data.A)  
6   jmeno zemepis  
7 1 Karel      20  
8 2 Pavel      10  
9 3  Jan       12  
10 4 Aneta      43  
11 5 Jakub      20
```

Pomocí argumentu `stringsAsFactors = FALSE` si zajistíme, že jména nebudou interpretována jako faktory. V případě, že by tomu tak bylo, tak bychom se mohli potýkat s problémy, protože faktor má charakter jen jako vizuální reprezentaci pořadové hodnoty faktoru, ale všechny výpočty jsou prováděné s touto pořadovou hodnotou. Druhý dat

obsahuje analogický seznam žáků a jejich počet bodů z matematiky.

```
1 > data.B <- data.frame(  
2 +       jmeno = c("Pavel", "Karel", "Jakub", "Aneta", "Pepa"),  
3 +       matematika = c(42, 55, 84, 87, 45),  
4 +       stringsAsFactors = FALSE)  
5 > print(data.B)  
6   jmeno matematika  
7 1 Pavel          42  
8 2 Karel          55  
9 3 Jakub          84  
10 4 Aneta          87  
11 5 Pepa           45
```

V případě, kdy na oba dva datasety aplikujeme prostou funkci `merge()`, tak sice dojde ke spojení datasetů podle předem zvoleného klíče, ale(!) o část dat přijdeme. Je to tím, že ve výchozím nastavení se na sebe napojí data, která jsou průsečíkem klíčů v obou datasetech. Za klíč zde považujeme proměnnou `jmeno`. Jak na to ale funkce přišla, že se jedná o klíč? Je to z toho důvodu, že se jedná o jedinou proměnnou, která má stejné jméno napříč oběma soubory a obsahuje charaktery.

```
1 > merge(x = data.A, y = data.B)  
2   jmeno zemepis  matematika  
3 1 Aneta       43          87  
4 2 Jakub       20          84  
5 3 Karel       20          55  
6 4 Pavel       10          42
```

Pokud bychom nastavili parametr `all` na hodnotu `TRUE`, získali bychom všechna data z obou souborů.

```
1 > merge(x = data.A, y = data.B, all = TRUE)  
2   jmeno zemepis  matematika  
3 1 Aneta       43          87  
4 2 Jakub       20          84  
5 3 Jan        12          NA  
6 4 Karel       20          55  
7 5 Pavel       10          42  
8 6 Pepa        NA          45
```

V případě chybějících hodnot proměnných funkce dosadí do výsledného datasetu `NA` hodnoty. V případě datasetu, kde jsou rozdílně pojmenované proměnné nesoucí klíč, tak přiřadíme jména parametrů `by.x` a `by.y`.

```

1 > merge(x = data.A, y = data.B, all = TRUE, by.x = "jmeno", by.y = "
  jmeno)
2   jmeno zemepis  matematika
3 1 Aneta      43         87
4 2 Jakub      20         84
5 3 Jan        12         NA
6 4 Karel      20         55
7 5 Pavel      10         42
8 6 Pepa       NA         45

```

V situaci, kdy se sice klíč jmenuje stejně, ale v datasetech existuje více stejně pojmenovaných proměnných, využijeme parametru `by`. V takové situaci by se u shodných jmen proměnných objevily přípony. Jméno přípony můžeme specifikovat v parametru `suffixes`. Pro ilustraci tohoto přístupu vydefinujeme nový dataset `data.C`, který bude obsahovat stejně jako `data.B` body z matematiky.

```

1 > data.C <- data.frame(
2 +   jmeno = c("Petr", "Samuel", "Jakub", "Aneta", "Pepa"),
3 +   matematika = c(47, 89, 81, 87, 45),
4 +   stringsAsFactors = FALSE
5 +   .... [TRUNCATED]
6
7 > merge(x = data.B, y = data.C,
8 +   all = TRUE, by = "jmeno",
9 +   suffixes = c(".prvnni", ".druhe")
10 + )
11   jmeno  matematika.prvnni  matematika.druhe
12 1 Aneta                  87                 87
13 2 Jakub                  84                 81
14 3 Karel                  55                 NA
15 4 Pavel                  42                 NA
16 5 Pepa                   45                 45
17 6 Petr                   NA                 47
18 7 Samuel                 NA                 89

```

Parametry `all.x` a `all.y` nám umožňují specifikovat, zda mají být z daného datasetu (přřazeného do proměnné `x` anebo `y`) přebrané všechny hodnoty, anebo jen ty hodnoty, které jsou průsečíkem na druhý dataset. Pokud nastavíme parametr `all` na hodnotu `TRUE`, tak automaticky oba parametry `all.x` a `all.y` přeberou hodnotu `TRUE`. Alternativně můžeme nastavit způsob průniku dat pro dílčí datasety zvlášť. Ukažme si postupně všechny varianty.

Tímto způsobem získáme jen data, která jsou podle klíče v obou datasetech.

```
1 > merge(x = data.A, y = data.B, all.x = FALSE, all.y = FALSE)
2   jmeno  zemepis  matematika
3 1 Aneta      43      87
4 2 Jakub     20      84
5 3 Karel     20      55
6 4 Pavel     10      42
```

V případě, kdybychom chtěli získat všechna data z datasetu `data.A` a zároveň k nim napárovat data z datasetu `data.B` (ale jen ta, co jsou obsažená v datasetu A), tak to provedeme pomocí následujícího příkazu.

```
1 > merge(x = data.A, y = data.B, all.x = TRUE, all.y = FALSE)
2   jmeno  zemepis  matematika
3 1 Aneta      43      87
4 2 Jakub     20      84
5 3 Jan       12      NA
6 4 Karel     20      55
7 5 Pavel     10      42
```

Opak — tedy pokud chceme získat všechna data z datasetu `data.B` a k nim přidat jen data, co mají stejný klíč z datasetu `data.A`, získáme je podle následujícího příkazu.

```
1 > merge(x = data.A, y = data.B, all.x = FALSE, all.y = TRUE)
2   jmeno  zemepis  matematika
3 1 Aneta      43      87
4 2 Jakub     20      84
5 3 Karel     20      55
6 4 Pavel     10      42
7 5 Pepa      NA      45
```

Analogii k parametru `all=TRUE` získáme nastavením obou hodnot na `TRUE`. Získáme tím dataset, kde budou obsažené všechny záznamy.

```
1 > merge(x = data.A, y = data.B, all.x = TRUE, all.y = TRUE)
2   jmeno  zemepis  matematika
3 1 Aneta      43      87
4 2 Jakub     20      84
5 3 Jan       12      NA
6 4 Karel     20      55
7 5 Pavel     10      42
8 6 Pepa      NA      45
```


7.2 Knihovna dplyr

Balíček `dplyr` obsahuje širokou škálu na práci s daty. Použití této knihovny je vhodné zejména v situaci, kdy máme velké datové soubory. Knihovna pracuje jak s objektem `data.frame`, tak ale zároveň navrácí objekt typu `tibble`. Objekt typu `tibble` je velmi podobný objektu typu `data.frame`. Prvním rozdílem je to, že pokud zavoláme jméno daného objektu v příkazové řádce, tak se nám automaticky vypíše jen hlavička souboru. Tedy provede se obdoba funkce `head()`. Dalším rozdílem je to, že v objektu typu `data.frame` jsou všechny proměnné typu `character` ve výchozím nastavení převáděné na `factor`, v objektu typu `tibble` zůstávají proměnné `characterem`.

7.2.1 Pipeline operátor

Knihovna `dplyr` zavádí nový operátor jménem pipeline (`%>%`). Ukážeme si nejprve na následujícím příkladu jeho použití na funkci `mean()`.

```
1 > znamky <- c(10, 20, 30)
2 > mean(znamky)
3 [1] 20
4 > znamky %>% mean
5 [1] 20
```

Standardní použití funkce `mean` spočívá v tom, že do kulaté závorky za jménem funkce vložíme do parametrů příslušné hodnoty. V tomto případě je ale objekt vlevo před operátorem `%>%` vložen jako hlavní (první) parametr do funkce napravo od operátoru. Nenechte se zmást tím, že byste si mysleli, že se jedná o porovnávání – vůbec nic zde neporovnáváme. Ukážeme si ještě jeden příklad s použitím funkce `colMeans()` na objektu `mtcars`.

```
1 > mtcars %>% colMeans
2      mpg      cyl      disp      hp
3 20.090625  6.187500 230.721875 146.687500
4      drat      wt      qsec      vs
5  3.596563  3.217250 17.848750  0.437500
6      am      gear      carb
7  0.406250  3.687500  2.812500
```

V situaci, kdy potřebujeme výsledek jedné funkce vložit jako vstup do jiné, nachází tento operátor použití. Takovouto situaci zřehlední a čtenář příslušného kódu ihned uvidí posloupnost operací. Ukážeme si to na následujícím příkladu výpočtu průměrů.

```
1 > mtcars %>% colMeans %>% mean
2 [1] 39.60853
```

```

3 > mean(colMeans(mtcars))
4 [1] 39.60853

```

7.2.2 Filtrování pozorování

Pro filtrování pozorování máme vícero možností. Nejprve můžeme filtrovat pomocí funkce `filter()`. Funkce `filter()` přebírá jako parametr logickou podmínku, kterou musí dané řádky splnit. Ukažme si to na datasetu `mtcars`.

```

1 > mtcars %>% filter(cyl == 6)
2   mpg  cyl  disp  hp drat    wt  qsec vs am
3 1 21.0    6 160.0 110 3.90 2.620 16.46 0  1
4 2 21.0    6 160.0 110 3.90 2.875 17.02 0  1
5 3 21.4    6 258.0 110 3.08 3.215 19.44 1  0
6 4 18.1    6 225.0 105 2.76 3.460 20.22 1  0
7 5 19.2    6 167.6 123 3.92 3.440 18.30 1  0
8 6 17.8    6 167.6 123 3.92 3.440 18.90 1  0
9 7 19.7    6 145.0 175 3.62 2.770 15.50 0  1
10
11 gear carb
12 4      4
13 4      4
14 3      1
15 3      1
16 4      4
17 4      4
18 5      6

```

Tímto způsobem jsme si vybrali řádky datasetu, které splňují podmínku, že mají 6 cylindrů. Podmínky můžeme i kombinovat pomocí standardních logických operátorů.

```

1 > mtcars %>% filter(cyl == 6 & disp > 170)
2   mpg  cyl  disp  hp drat    wt  qsec vs am
3 1 21.4    6 258 110 3.08 3.215 19.44 1  0
4 2 18.1    6 225 105 2.76 3.460 20.22 1  0
5   gear carb
6 1     3    1
7 2     3    1

```

Funkce `distinct` nám z datasetu odstraní duplicitní záznamy. Ukažme si to na následujícím příkladu.

```

1 > data.X <- data.frame(x1 = c(1, 1, 2, 1, 2),
2 +                       x2 = c(1, 1, 2, 3, 2))
3 > print(data.X)
4   x1 x2

```

```

5 1 1 1
6 2 1 1
7 3 2 2
8 4 1 3
9 5 2 2
10 > data.X %>% distinct()
11   x1 x2
12 1 1 1
13 2 2 2
14 3 1 3

```

Náhodný výběr pozorování lze provést pomocí dvou funkcí. První funkce je `sample_frac()`, která definuje podíl (v procentním čísle) dat, který se má vybrat z datasetu, druhá funkce jménem `sample_n()` definuje počet řádek, který se má vybrat z datasetu.

```

1 > iris %>% sample_n(2)
2   Sepal.Length Sepal.Width Petal.Length Petal.Width   Species
3 1           7.2         3.0          5.8         1.6 virginica
4 2           5.4         3.7          1.5         0.2   setosa
5 > iris %>% sample_frac(0.025)
6   Sepal.Length Sepal.Width Petal.Length Petal.Width   Species
7 1           5.1         3.5          1.4         0.2   setosa
8 2           6.9         3.1          5.1         2.3 virginica
9 3           5.2         3.4          1.4         0.2   setosa
10 4           5.1         3.3          1.7         0.5   setosa

```

Filtrování funkcí můžeme provést i pomocí kombinace s matematickými funkcemi.

```

1 > mtcars %>% filter(hp > mean(hp) & wt < mean(wt))
2   mpg cyl disp  hp drat   wt  qsec vs am gear carb
3 1 15.8   8  351 264 4.22 3.17 14.5  0  1    5    4
4 2 19.7   6  145 175 3.62 2.77 15.5  0  1    5    6

```

7.2.3 Setřídění datasetu

Dataset můžeme snadno i setřídít. Stačí použít funkci `arrange()`. Funkce `arrange()` ve výchozím nastavení třídí data vzestupně.

```

1 > setridene = mtcars %>% arrange(hp)
2 > head(setridene)
3   mpg cyl  disp hp drat   wt  qsec vs am gear carb
4 1 30.4   4  75.7 52 4.93 1.615 18.52  1  1    4    2
5 2 24.4   4 146.7 62 3.69 3.190 20.00  1  0    4    2
6 3 33.9   4  71.1 65 4.22 1.835 19.90  1  1    4    1
7 4 32.4   4  78.7 66 4.08 2.200 19.47  1  1    4    1

```

```

8 5 27.3 4 79.0 66 4.08 1.935 18.90 1 1 4 1
9 6 26.0 4 120.3 91 4.43 2.140 16.70 0 1 5 2

```

Pro sestupní třídění je nutné proměnnou uvnitř funkce `arrange()` vložit do funkce `desc` – descending.

```

1 > setridene = mtcars %>% arrange(desc(hp))
2 > head(setridene)
3   mpg  cyl  disp  hp  drat    wt   qsec  vs  am  gear  carb
4 1 15.0   8  301 335 3.54 3.570 14.60 0  1    5    8
5 2 15.8   8  351 264 4.22 3.170 14.50 0  1    5    4
6 3 14.3   8  360 245 3.21 3.570 15.84 0  0    3    4
7 4 13.3   8  350 245 3.73 3.840 15.41 0  0    3    4
8 5 14.7   8  440 230 3.23 5.345 17.42 0  0    3    4
9 6 10.4   8  460 215 3.00 5.424 17.82 0  0    3    4

```

Data můžeme třídit i podle více proměnných, stačí je vložit za sebou do funkce `arrange()`.

```

1 > setridene = mtcars %>% arrange(gear, carb)
2 > head(setridene)
3   mpg  cyl  disp  hp  drat    wt   qsec  vs  am  gear  carb
4 1 21.4   6 258.0 110 3.08 3.215 19.44 1  0    3    1
5 2 18.1   6 225.0 105 2.76 3.460 20.22 1  0    3    1
6 3 21.5   4 120.1  97 3.70 2.465 20.01 1  0    3    1
7 4 18.7   8 360.0 175 3.15 3.440 17.02 0  0    3    2
8 5 15.5   8 318.0 150 2.76 3.520 16.87 0  0    3    2
9 6 15.2   8 304.0 150 3.15 3.435 17.30 0  0    3    2

```

Pořadí proměnných ve funkci `arrange()` ovlivňuje i prioritu, podle které jsou výsledná data setříděná. Prvně jsou data setříděna podle skupin proměnné `gear` a následně je v těchto skupinách je pořadí řízené pomocí proměnné `carb`.

7.2.4 Filtrování proměnných

Funkce `select()` nám umožní filtrovat dataset z pohledu proměnných. Z datasetu `mtcars` provedeme selekci výběr `vs` a `am`, dále si zobrazíme hlavičku souboru.

```

1 > selekt = mtcars %>% select(vs, am)
2 > head(selekt)
3           vs  am
4 Mazda RX4      0  1
5 Mazda RX4 Wag  0  1
6 Datsun 710     1  1
7 Hornet 4 Drive  1  0
8 Hornet Sportabout 0  0
9 Valiant        1  0

```

Do funkce `select()` stačí vypsát jména proměnných, které chceme z datasetu vybrat. Pokud bychom chtěli vybrat všechna data, tak jako parametr vložíme funkci `everything()`.

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

Funkce `everything()` má zejména použití v dalších funkcích, které plní určitou úlohu, ale zároveň provádí i výběr proměnných. Popřípadě v situaci, kdy chceme přetřídit proměnné datasetu. Funkce `everything()` pak sama dovybere zbývající proměnné.

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

Pokud bychom chtěli vybrat proměnné, které obsahují ve svém jménu specifický řetězec, tak použijeme funkci `contains()`.

```
1 > selekt = mtcars %>% select(contains("a"))
2 > head(selekt)
3           drat  am  gear  carb
4 Mazda RX4     3.90   1    4    4
```

5	Mazda RX4 Wag	3.90	1	4	4
6	Datsun 710	3.85	1	4	1
7	Hornet 4 Drive	3.08	0	3	1
8	Hornet Sportabout	3.15	0	3	2
9	Valiant	2.76	0	3	1

Alternativně můžeme vybírat i podle konců a začátků jmen. Pomocí funkcí `starts_with()` a `ends_with()`.

```

1 > selekt = mtcars %>% select(ends_with("p"))
2 > head(selekt)
3
4           disp  hp
5 Mazda RX4    160 110
6 Mazda RX4 Wag 160 110
7 Datsun 710   108  93
8 Hornet 4 Drive 258 110
9 Hornet Sportabout 360 175
10 Valiant     225 105

```

7.2.5 Tvorba nové proměnné

Novou proměnnou vytvoříme velmi snadno. Zdefinujeme ji uvnitř funkce `mutate()`. V následujícím příkladu normujeme pomocí z-skóre proměnnou `mpg`, kterou si uložíme nově jako „Z_mpg“.

```

1 > data.x <- mtcars
2 > data.x <- data.x %>% mutate(Z_mpg = (mpg - mean(mpg))/sd(mpg))
3 > head(data.x)
4   mpg cyl disp  hp drat   wt  qsec vs am gear carb  Z_mpg
5 1 21.0   6  160 110 3.90 2.620 16.46  0  1   4    4  0.1508848
6 2 21.0   6  160 110 3.90 2.875 17.02  0  1   4    4  0.1508848
7 3 22.8   4  108  93 3.85 2.320 18.61  1  1   4    1  0.4495434
8 4 21.4   6  258 110 3.08 3.215 19.44  1  0   3    1  0.2172534
9 5 18.7   8  360 175 3.15 3.440 17.02  0  0   3    2 -0.2307345
10 6 18.1   6  225 105 2.76 3.460 20.22  1  0   3    1 -0.3302874

```

Pomocí funkce `mutate_each()` můžeme přepočítat všechny proměnné. Tato funkce přebírá jako parametr jméno funkce, která se má na dané sloupce použít, popřípadě deklaraci takové funkce. V našem případě od proměnné odečteme průměr a zaokrouhlíme hodnoty na celá čísla.

```

1 > data.x <- mtcars
2 > data.x <- data.x %>% mutate_each(function(x) round(x - mean(x)))
3 > head(data.x)
4   mpg cyl disp  hp drat wt  qsec vs am gear carb

```

5	1	1	0	-71	-37	0	-1	-1	0	1	0	1
6	2	1	0	-71	-37	0	0	-1	0	1	0	1
7	3	3	-2	-123	-54	0	-1	1	1	1	0	-2
8	4	1	0	27	-37	-1	0	2	1	0	-1	-2
9	5	-1	2	129	28	0	0	-1	0	0	-1	-1
10	6	-2	0	-6	-42	-1	0	2	1	0	-1	-2

Funkce `transmute()` pak provádí změnu jen jedné konkrétní proměnné. Díky ní můžeme transformovat/přepočítat jednu či více zvolených proměnných. V následujícím příkladu zaokrouhlíme proměnnou `mpg` na celá čísla. Důležité je si uvědomit, že funkce `transmute()` funguje zároveň jako funkce `selekt()` a vybírá sloupce datasetu.

```

1 > data.x <- mtcars
2 > data.x <- data.x %>% transmute(mpg = round(mpg))
3 > head(data.x)
4   mpg
5 1  21
6 2  21
7 3  23
8 4  21
9 5  19
10 6  18

```

7.2.6 Statistické shrnutí dat

Knihovna `dplyr` disponuje i možností tvorby základních statistik pro daný dataset. Jednu proměnnou z datasetu můžeme shrnout pomocí funkce `summarise()`. Vstupem do této funkce je deklarace proměnné provádějící statistické shrnutí. Dostupné funkce pro statistické shrnutí jsou: `first()` – první hodnota v datasetu, `last()` – poslední hodnota, `min()` – minimální hodnota, `max()` – maximální hodnota, `nth()` – n-tá hodnota datasetu (např. pořadové číslo řádky), `mean()` – průměrná hodnota, `n()` – počet hodnot z datasetu, `median()` – mediánová hodnota, `var()` – rozptyl, `sd()` – směrodatná odchylka, `IQR()` – mezikvartilové rozpětí.

```

1 > data.x <- mtcars
2 > data.x %>% summarise(wt_mean = mean(wt))
3   wt_mean
4 1 3.21725

```

Pomocí funkce `summarise_all` můžeme shrnout i celý dataset. Ukažme si alternativu funkce `colMeans()` v následujícím příkladu.

```

1 > data.x %>% summarise_all(mean)
2   mpg      cyl    disp      hp      drat
3 1 20.09062  6.1875 230.7219 146.6875  3.596563

```

	wt	qsec	vs	am	gear	carb
1	3.21725	17.84875	0.4375	0.40625	3.6875	2.8125

7.2.7 Práce se skupinami dat

Princip práce se skupinami dat v knihovně dplyr spočívá v tom, že každému řádku v datasetu přiřadíme identifikaci k určité skupině. Následně všechny operace provedené na celém datasetu jsou prováděné podle těchto skupin zvlášť. Funkce, která slouží k rozdělení dat do skupin, se jmenuje `group_by()`. Zkusme si tuto úlohu provést pro výpočet průměrů v datasetu `mtcars`, který roztřídíme do skupin podle válců.

```
1 > data.x %>% group_by(cyl) %>% summarise_all(mean)
2 # A tibble: 3 x 11
3   cyl   mpg  disp   hp  drat   wt   qsec
4   <dbl> <dbl> <dbl> <dbl> <dbl> <dbl> <dbl>
5 1     4  26.7  105.  82.6  4.07  2.29  19.1
6 2     6  19.7  183.  122.   3.59  3.12  18.0
7 3     8  15.1  353.  209.   3.23  4.00  16.8
8 # ... with 4 more variables: vs <dbl>,
9 #   am <dbl>, gear <dbl>, carb <dbl>
```

Tato úloha je velmi podobná funkci `aggregate()` s parametrem `by` nastaveným na počet válců.

7.2.8 Spojování datasetů pomocí balíčku dplyr

Knihovna `dplyr` disponuje vlastními funkcemi na spojení datasetů. Definujme si stejná data jako v případě funkce `merge()`.

```
1 > data.A <- data.frame(
2 +   jmeno = c("Karel", "Pavel", "Jan", "Aneta", "Jakub"),
3 +   zemepis = c(20, 10, 12, 43, 20),
4 +   stringsAsFactors = FALSE
5 + )
6 > data.B <- data.frame(
7 +   jmeno = c("Pavel", "Karel", "Jakub", "Aneta", "Pepa"),
8 +   matematika = c(42, 55, 84, 87, 45),
9 +   stringsAsFactors = FALSE
10 + )
11 > print(data.A)
12   jmeno zemepis
13 1 Karel      20
14 2 Pavel      10
15 3 Jan       12
```



```

16 4 Aneta      43
17 5 Jakub      20
18 > print(data.B)
19   jmeno matematika
20 1 Pavel      42
21 2 Karel      55
22 3 Jakub      84
23 4 Aneta      87
24 5 Pepa       45

```

Funkce `left_join()` provede spojení všech dat z datasetu `data.A` a přiřadí k nim data z datasetu `data.B`.

```

1 > data.A %>% left_join(data.B, by = "jmeno")
2   jmeno zemepis matematika
3 1 Karel      20         55
4 2 Pavel      10         42
5 3 Jan        12         NA
6 4 Aneta      43         87
7 5 Jakub      20         84

```

Funkce `right_join()` naopak provede spojení všech dat z datasetu `data.B` a k nim přiřadí data z datasetu `data.A`.

```

1 > data.A %>% right_join(data.B, by = "jmeno")
2   jmeno zemepis matematika
3 1 Pavel      10         42
4 2 Karel      20         55
5 3 Jakub      20         84
6 4 Aneta      43         87
7 5 Pepa       NA         45

```

Funkce `inner_join()` spojí data podle klíče, který se vyskytuje v obou datasetech.

```

1 > data.A %>% inner_join(data.B, by = "jmeno")
2   jmeno zemepis matematika
3 1 Karel      20         55
4 2 Pavel      10         42
5 3 Aneta      43         87
6 4 Jakub      20         84

```

Pokud chceme spojit data z obou datasetů a zachovat všechny hodnoty (obdoba parametru `all=TRUE` u funkce `merge`), použijeme funkci `full_join()`.

```

1 > data.A %>% full_join(data.B, by = "jmeno")
2   jmeno zemepis matematika
3 1 Karel      20         55

```

4	2	Pavel	10	42
5	3	Jan	12	NA
6	4	Aneta	43	87
7	5	Jakub	20	84
8	6	Pepa	NA	45

Kapitola 8

Formátování kódu

Při psaní kódu je vhodné dodržet základní pravidla formátování. Existuje vícero druhů pravidel, jak formátovat kód v jazyce R. Zde si představíme formát sestavený Hadley Wickhamem (Wickham, H. 2020) v rámci knihovny tidyverse ([odkaz zde](#)). Další formáty jsou například od Jean Fan (Fan, J., 2020; [odkaz zde](#)) či Google community style guide ([odkaz zde](#)). Časem jistě zjistíte, že existuje poměrně mnoho formátů. Čtenáři doporučujeme vybrat si jeden formát, a pokud je to možné, tak se ho dále držet. Reformátování vašeho kódu můžete v RStudios provést velmi snadno pomocí tlačítka na reformátování kódu. Standard zde použitý odpovídá standardu Hadley Wickhama Code -> Reformat Code (či CTRL+SHIFT+A).

Proč ale kód specificky formátovat? Formátování kódu má mnoho výhod. Za hlavní výhodu se obecně považuje snadnější srozumitelnost a čitelnost. Zde bychom rádi čtenáře varovali, aby si nemyslel, že když nějaký kód naprogramuje, tak mu bude vždy rozumět. Pokud jste nedodrželi rozumné zásady formátování, tak zjistíte, že je vám váš vlastní kód nesrozumitelný.

Pokud je to možné, tak je vhodné dodržet zásady formátování všude, kde to jde. Dále by se uživatel měl snažit psát srozumitelný kód a volit postupy, co jsou srozumitelné. Může se vám ale stát, že snaha psát srozumitelný a přehledný kód narazí na jeho výkonnost. Tedy, že jedna věc napsána přehledně nemusí být napsána výkonně z pohledu rychlosti běhu samotného kódu. Je tedy na uživateli, aby zvážil tyto všechny aspekty programování: formátovací styl, použité funkce a algoritmy a srozumitelnost.

V následujících pár řádcích se vám pokusíme nastínit základy formátování představené již zmíněným Hadleyem Wickhamem.

8.1 Jména

8.1.1 Jména objektů a funkcí

V případě jmen objektů a funkcí volíme vhodné pojmenování, které vystihuje daný jev. Je správné volit krátké, ale výstižné názvy. Pro oddělování slov v názvech je vhodné použít podtržítko (`_`) a vyvarovat se tzv. maďarské notace, kdy jsou slova spojována k sobě bez oddělovacího znaku a rozlišena malým/velkým písmenem. Použití tečky pro oddělení slov není vyloženě špatné, ale ani doporučované. Ukážeme si na názvu, který má být pro objekt, co ponese průměr z prvočísel.

```
1 # Dobře
2 prumer_prvocisla
3
4
5 # Ani dobře ani špatně:
6 prumer.prvocisla
7 prum_prvoc
8
9 # Špatně
10 PrumerPrvocisla
11 prumerPrvocisla
12 prumerprvocisla
13 prumer_vypocitany_z_mejch_prvocisel
14 pp
15 p
```

8.1.2 Jména souborů

Stejná pravidla platí i pro pojmenování souborů, s tou výjimkou, že v názvu souboru můžeme na prvním místě použít i čísla a nově i další speciální znak pomlčku (`-`). Vyvarujte se použití mezery a tečky v názvu souboru (netýká se přípony).

```
1 # Dobře
2 "01-uvod.R"
3 "02-funkce.R"
4 "03-analyza_dat.R"
5 ...
6 "10-export_dat.R"
7
8 # Špatně
9 "01 uvod.R"
10 "2.funkce.R"
11 "3-analyzaDat.R"
```

V případě, kdy soubory číslyjete, je vhodné dodržet délku znaků odpovídající maximální hodnotě čísla.

8.2 Obecná syntaxe

8.2.1 Psaní mezer

V případě výběrů a parametrů funkcí byste měli za čárkou psát mezeru.

```
1 # Dobře
2 X[1, 2]
3 X[1, 2, 45]
4 X[, 3]
5 mean(x, na.rm = TRUE)
6
7 # Spatne
8 X[1,2]
9 X[1,]
10 X[ , , 3]
11 mean(x,na.rm = TRUE)
```

Toto pravidlo platí i pro všechny operátory. Za a po operátoru byste měli psát mezeru. Na začátku závorek a na konci se mezera nepíše a mezi jménem funkce a závorkou se mezera nepíše také!

```
1 # Dobře
2 X %% Y
3 x <- matrix(1:9, ncol = 3)
4
5 # Spatne
6 X%%Y
7 x<-matrix(1:9,ncol=3)
8 x<-matrix( 1:9, ncol = 3 )
9 x <- matrix (1:9, ncol = 3)
```

Výjimka je v případě následujících operátorů.

```
1 :, ::, :::, $, @, [, [[ a ^.
```

```
1 # Dobře
2 1:10
3 fit$coef
4 X[, 3]
5 X[[1]]
6
```

```

7 # Spatne
8 1 : 10
9 fit $ coef
10 X [, 3]
11 X [[1]]

```

V případě funkcí, kdy vynecháváme nějaký argument, tak nevytváříme prázdné místo pro tento argument pomocí čárek.

```

1 # Dobře
2 seq(from = 1, to = 10, length.out = 5)
3
4 # Spatne
5 seq(from = 1, to = 10, , length.out = 5)

```

V případě cyklů a podmínek píšeme mezeru mezi jménem cyklu a závorkou. Tato mezera se týká i případných složených závorek.

```

1 # Dobře
2 for (i in 1:10) {
3   print(i)
4 }
5
6 if (i > 4) {
7   print(i)
8 }
9 # Spatne
10 for(i in 1:10) {
11   print(i)
12 }
13
14 if (i > 4){
15   print(i)
16 }
17 if(i > 4){
18   print(i)
19 }

```

8.2.2 Logické operátory

U logických operátorů se vyvarujte použití zkrácených názvů, kterými jsou T a F, ale používejte plné termíny.

```

1 # Dobře
2 TRUE

```

```
3 FALSE
4
5 # Spatne
6 T
7 F
```

8.2.3 Uvozovky

V případě používání uvozovek se doporučuje vyvarovat se použití znaku ' a místo toho používat jen znak ".

```
1 # Dobre
2 print("nas text")
3
4 # Spatne
5 print('nas text')
```

Pozor, tyto dva znaky nejsou zcela zastupitelné, proto existují situace, kdy musíte znak ' použít.

8.3 Bloky kódu a delší zápis kódu

Bloky kódu, které jsou dané do složených závorek, by měly být osazené tabulátorem. Toto se týká zejména cyklů, podmínek a uživatelem definovaných funkcí. Zároveň mějme na paměti doporučení o tom, že řádek kódu by neměl přesáhnout 80 znaků.

```
1 # Dobre
2 for (i in 1:10) {
3   print(i)
4   if (i > 5){
5     print("je vetsi")
6   }
7 }
8
9 moje_funkce <- function(X) {
10   print(X)
11 }
12
13 moje_funkce <- function(X = "vychozi argument 1",
14                           Y = "vychozi argument 2") {
15   print(X)
16 }
17
```

```

18 # Spatne
19
20 for (i in 1:10) {
21   print(i)
22   if (i > 5){
23     print("je vetsi")
24   }
25 }
26
27
28 moje_funkce <- function(X="vychozi argument 1", Y="vychozi argument
    2") {
29   print(X)
30 }
31
32 moje_funkce <-function(X = "vychozi argument 1",
33   Y = "vychozi argument 2") {
34   print(X)
35 }
36
37 moje_funkce <- function(X = "vychozi argument 1",
38   Y = "vychozi argument 2") {
39   print(X)
40 }

```

Formátování kódu je velmi důležité i pro spolupráci s ostatními programátory a dobře formátovaným kódem mnoho lidí potěšíte. Naopak špatně formátovaným kódem můžete mnoho lidí rozčítit. Časem si na formátování kódu zvyknete a bude mít pro vás stejný význam jako gramatika v českém jazyce.

Seznam obrázků

4.1	Plot funkce	61
6.1	Funkce plot()	81
6.2	Funkce plot() – úprava popisků	82
6.3	Základních 657 barev jazyka R	85
6.4	Funkce plot() – úprava parametrů	87
6.5	Graf sinusoidy	89
6.6	Funkce plot() a více proměnných	91
6.7	Funkce plot() – ylim a xlim parametry	92
6.8	Funkce barplot()	93
6.9	Úprava parametrů funkce barplot()	94
6.10	Funkce boxplot()	95
6.11	Funkce boxplot() – úprava parametrů funkce	97
6.12	Funkce hist()	98
6.13	Funkce hist() a hist(freq=FALSE)	99
6.14	Funkce hist() – úprava parametrů funkce	100
6.15	Funkce dotchart()	101
6.16	Funkce dotchart() – úprava parametrů funkce	103
6.17	Funkce pie()	104
6.18	Funkce stripchart()	105
6.19	Funkce stripchart() se specifikací jitter	106
6.20	Více grafů v jedné grafické oblasti funkce par()	108
6.21	Layout model grafu	109
6.22	Layout model grafu II	109
6.23	Více grafů v jedné grafické oblasti layout – funkce	111
6.24	Okraje grafu – outer a inner okraj	113
6.25	Úprava okrajů grafu	115
6.26	Tvorba legendy grafu	117
6.27	Tvorba legendy okrajů grafu mimo oblast grafu	119
6.28	Ukládání grafů	121
6.29	Funkce ggplot()	125
6.30	Funkce ggplot: geom_point	126
6.31	Funkce ggplot: geom_histogram	128

6.32	Funkce ggplot: geom_point, geom_smooth	129
6.33	Funkce ggplot: geom_point, geom_smooth pro data rozdělená do skupin	130
6.34	Funkce ggplot: geom_point, geom_smooth pro data celkem	131
6.35	Funkce ggplot: geom_point, geom_smooth pro skupiny	132
6.36	Funkce ggplot: geom_line	133
6.37	Funkce ggplot: geom_bar	134
6.38	Funkce ggplot: geom_histogram	135
6.39	Funkce ggplot: scale funkce I	136
6.40	Funkce ggplot: scale funkce II	137
6.41	Funkce ggplot: scale funkce III	138
6.42	Funkce ggplot: scale funkce IV	139
6.43	Funkce ggplot: ggtitle, xlab a ylab	140
6.44	Funkce ggplot: labs	141
6.45	Funkce ggplot: theme	143
6.46	Funkce ggplot: facet_grid funkce I	144
6.47	Funkce ggplot: facet_grid funkce II	145
6.48	Funkce ggplot: facet_grid funkce, scales=free	146
6.49	Funkce ggplot: facet_wrap funkce	147

Seznam tabulek

6.1 Parametry funkcí 83

6.2 Druhy čar 86

Literatura

- [1] FAN, J. *R Style Guide: Best practices for readable, sharable, and verifiable R code*[online]. JEFworks 2015-2020, 2020 [cit. 2020-02-02]. Dostupné z: <http://jef.works/R-style-guide/>
- [2] GOOGLE. *Google's R Style Guide* [online]. Google Style Guides, 2020 [cit. 2020-02-05]. Dostupné z: <https://google.github.io/styleguide/Rguide.html>
- [3] LONG, J. D., TEETOR, P., 2019. *R Cookbook: Proven Recipes for Data Analysis, Statistics, and Graphics*. O'Reilly Media, 2. vydání. ISBN: 1492040681.
- [4] RAHLF T., 2019. *Data Visualisation with R – 111 Examples*, Cham: Springer Nature, 2 vydání. ISBN: 978-3-030-28444-2.
- [5] R CORE TEAM. *A language and environment for statistical computing. R Foundation for Statistical Computing* [online]. Vienna, Austria: R Foundation, 2017, [cit. 2020-01-17]. Dostupné z: <https://www.R-project.org/>
- [6] R CORE TEAM. *R: A language and environment for statistical computing. R Foundation for Statistical Computing*[online]. Vienna, Austria: R Foundation, 2020 [cit. 2020-01-20]. Dostupné z: <https://www.R-project.org/>
- [7] RStudio Team. *RStudio: Integrated Development for R* [online]. Boston, MA: RStudio, Inc., 2018 [cit. 2020-01-12]. Dostupné z: <http://www.rstudio.com/>
- [8] SPECTOR, P., 2008. *Data manipulation with R*. New York: Springer. Use R! ISBN 978-0-387-74730-9.
- [9] VENABLES, W. N. SMITH, D. M., R Core Team. *An Introduction to R. Notes on R: A Programming Environment for Data Analysis and Graphics 2020* [online]. Cran R: The Comprehensive R Archive Network, 2020 [cit. 2020-04-14]. Dostupné z: <https://cran.r-project.org/doc/manuals/r-release/R-intro.pdf>
- [10] WICKHAM, H., 2009. *Ggplot2: elegant graphics for data analysis*. Dordrecht: Springer. Use R! (Springer). ISBN 978-0-387-98140-6.

-
- [11] WICKHAM, H. *Advanced R* [online]. Hadley Wickham, 2020a [cit. 2020-04-16]. Dostupné z: <http://adv-r.had.co.nz/>
- [12] WICKHAM, H.. *The tidyverse style guide* [online]. Tidiverse 2020b [cit. 2020-02-01]. Dostupné z: https://style.tidyverse.org/_main.pdf
- [13] WICKHAM H., ROMAIN F., LIONEL H., MULLER, K., RStudio. *dplyr: A Grammar of Data Manipulation* [online]. Cran R: The Comprehensive R Archive Network, 2020 [cit. 2020-04-16]. Dostupné z: <https://cran.r-project.org/web/packages/dplyr/index.html>

Název	R snadno a rychle 2 Vizualizace dat a programování
Autoři	Ing. Karel Šafr, Ph.D. Ing. Jakub Danko, PhD.
Vydavatel	Vysoká škola ekonomická v Praze Nakladatelství Oeconomica
Vydání	první vydání v elektronické podobě
Redakční úprava	Mgr. Ludmila Doudová
Grafický návrh	Daniel Hamerník, DiS.
Počet stran	174
DTP	Vysoká škola ekonomická v Praze Nakladatelství Oeconomica
Doporučená cena	zdarma

ISBN 978-80-245-2381-1
