

Data1: Naměřené hodnoty studované veličiny, např. výška rostliny.

Data2: Zdvojený soubor Data1 (= každá hodnota z Data1 byla zkopírována).

Data1	Data2
1	1
5	5
7	7
8	8
9	9
10	10
11	11
12	12
12.5	12.5
13	13
14	14
15	15
16	16
17	17
18	18
20	20
24	24
	1
	5
	7
	8
	9
	10
	11
	12
	12.5
	13
	14
	15
	16
	17
	18
	20
	24

Příklad 2: Vlastnosti směrodatné odchylky a střední chyby průměru.

Použité proměnné: Data1 vs. Data2

Chování směrodatné odchylky a střední chyby průměru v závislosti na počtu měření si ukážeme na proměnné Data2, což je zdvojený soubor Data1. Aritmetický průměr souboru Data2 zůstal pochopitelně nezměněn. Vidíme také, že ani směrodatná odchylka souboru se (prakticky) nezměnila – zdvojením dat se hodnota čitatele při výpočtu směrodatné odchylky zdvojnásobila a hodnota jmenovatele skoro taktéž, pouze místo $(n-1)=16$ v prvním případě nyní máme $(n-1)=33$. Odměnou za dvojnásobný počet měření je ale podstatně vyšší míra jistoty odhadu parametrického průměru studované populace, tj. střední chyba průměru, která se snížila z 1.383 na 0.963. Je třeba si připomenout, že střední chyba průměru je vlastně směrodatná odchylka rozdělení výběrových průměrů, čili míra toho, jak se výběrové průměry liší mezi jednotlivými (hypotetickými) výběry.

Summary Section of Data1

Count	Mean	Standard Deviation	Standard Error	Minimum	Maximum	Range
17	12.5	5.700877	1.382666	1	24	23

Summary Section of Data2

Count	Mean	Standard Deviation	Standard Error	Minimum	Maximum	Range
34	12.5	5.613836	0.9627649	1	24	23