

K pojmu „repräsentativita“ v sociologických výzkumech

JAN ŘEHÁK

Ústav pro filozofii a sociologii ČSAV, Praha

Pojem „repräsentativita“, „repräsentativnost“, „repräsentativní šetření“ se velmi běžně vyskytuje v požadavech výzkumných projektů. Není však nijak specifikován a jeho použití je většinou nejasné, chybí mu konkrétní náplň. Často je chápán zcela chybně, ale především je ve většině případů velmi zjednodušován. Jeho použití je zpravidla jen intuitivní. Proto můžeme základní otázky pro tuto stat položit takto:

1. Co je to repräsentativita?
2. Co chceme repräsentovat?
3. Proč chceme repräsentativní šetření?
4. Jak můžeme repräsentativitu měřit a ověřovat?

Půjde zde o repräsentativitu spojenou s výběrovými šetřeními, to jest o to, co jsou a jaké mají být repräsentativní výběry a jakými postupy je v praxi zajišťujeme.

1. Vztah základního (cílového) souboru a výběrového souboru

Ve statistické analýze se snažíme oprostit získaná výzkumná data od různých chyb a zobecnit závěry z nich natolik, aby charakterizovaly předem určený základní soubor, který je předmětem a cílem našeho zájmu, a aby se staly základem pro vytváření sociologické informace, základem pro praktická rozhodování i podnětem pro sociologickou imaginaci.

Data jsou vždy zatížena velkým množstvím nejrůznějších chyb, které ztěžují výzkumný úkol; tyto chyby (měření, přenosu dat, jejich zpracování a sumarizace) jsou však zcela běžné a zákonité a nutně s nimi musíme v analýze a zobecňování počítat. Abstrahujeme-li od chyb měření, přenosu dat, zpracování i chyb sumarizace dat, je základním požadavkem výzkumu, aby data „repräsentovala“ určitý základní soubor, aby bylo možné přenášet vlastnosti výběrových dat na celou populaci.

Cíl výzkumných sociologických úvah má především dvě základní složky: obsah o němž pojednáváme (téma) a soubor, jehož se naše úsilí týká. V této stati pojednáme především o druhém aspektu, budeme hovořit o *cílové populaci*.

Výběrový soubor, to jest soubor dat, s nímž konkrétně pracujeme, zastupuje soubor základní, je jeho repräsentantem: repräsentativita může být tedy chápána ve smyslu zastupitelnosti cílové populace souborem výběrovým. Žádné výběrové šetření nemůže být přesnou zmenšeninou definované populace. Člověk, jeho prostředí, jeho situace, jeho zapojení do pracovních i jiných skupin, jeho vztahy k lidem i k jevům tvoří příliš variabilní, příliš rozmanité a heterogenní soubory, než aby mohla jedna osoba zastupovat druhou, jedna skupina zastupovat jinou skupinu a podobně. Člověk, pracovní skupina, rodina, třída žáků a podobně jsou prostě neopakovatelné sociologické objekty. V sociologických šetřeních se ovšem nezabýváme zcela detailním popisem či měřením objektů výzkumu. Provádíme redukci. Zcela detailní zachycení vlastností není ani vždy nutné, ani není dost dobře možné. A tak vzniká otázka

representativity obsahové: do jaké míry je v projektu zachycen obsah cílových témat. Sem patří problematika indikací a vhodných dekompozic pojmů, patří sem také stupeň pokrytí obsahu výzkumnými instrumenty, a to jak co do jeho šíře, tak do hloubky. Patří sem také validita a reliabilita údajů, tj. skutečná vypovídací hodnota dat. I když se tímto aspektem v práci hlouběji nezabýváme, nutno jej považovat za stejně důležitý jako reprezentativitu vztahenou k cílové populaci a implicitní našim dalším úvahám.

Chápeme-li úlohu zobecňování zcela pragmaticky, tj. ve vztahu k určitému, jednoznačně formulovanému sociologickému úkolu, pak ovšem nemusíme mluvit o reprezentativitě v plné šíři, ale o reprezentativitě výběrového souboru *vzhledem k danému úkolu*. Chceme-li zkoumat například názory čtenářů na změnu struktury časopisu, nemusíme přesně reprezentovat populaci čtenářů, ale stačí nám reprezentace určité části této populace, například těch čtenářů, kteří mají o časopis hlubší zájem. Chceme-li ovšem zkoumat efektivitu působení časopisu či znát postoje k časopisu v celé její heterogenitě u čtenářské populace jakožto celku, musíme zajistit reprezentativitu celé čtenářské obce.

Jednoduchý příklad z výzkumné praxe: Projekt *Automatizace a průmysloví dělníci* byl plánován tak, že bylo vybráno několik závodů v rámci automobilového průmyslu. V každém závodě byla vybrána jedna automatizovaná a jedna neautomatizovaná dílna (v poměru počtu respondentů 1 : 2 v každém závodě). Úkolem bylo komparovat dvě úrovně automatizovanosti výrobního procesu s přihlédnutím k heterogenitě dané specifickým vlivem lokality a závodu. Přesný postup výběru plně odpovídal této úloze. Výsledná data však nemohou reprezentovat ani závod či podnik, ani automobilový průmysl, ani dělnickou třídu celé země. Nemohou reprezentovat ani vývoj automatizace jako celek, neboť při omezení na dvě jednoznačné úrovně nelze podchytit celý proces zejména proto, že v jiných typech automatizace může docházet k jiným kvalitativním proměnám obsahu práce.⁽¹⁾ Reprezentativitu zde musíme chápat účelově, musíme ji vztáhnout k danému komparativnímu cíli.

Reprezentativita se vztahuje také k *pokrytí výzkumného instrumentu*. Rozlišujeme situace, v nichž reprezentujeme buď jednu nebo několik málo zkoumaných proměnných, nebo většinu proměnných, či všechny proměnné. Ne všechny znaky musí odrážet reálnou situaci cílové populace. Je zřejmé, že reprezentativitu máme zajištěnou především u těch znaků, u nichž si ji předepíšeme postupem: stratifikace, určení kvót. U jiných znaků se mohou projevit vlivy, jež reprezentativnost porušují. Týká se to především otázek, které jsou respondentům nepříjemné, týkají se intimních vztahů a problémů, přestupků proti předpisům, trestů, věcí, za něž se respondenti stydí, názorů na spolupracovníky apod. U těchto a mnoha jiných proměnných bývá reprezentativita porušena neochotou⁽²⁾ odpovídat, záměrným zkreslováním odpovědí či přímým odmítnutím.

Reprezentativnost může být odstupňována nejen vzhledem k počtu znaků obsažených ve výzkumných instrumentech, ale také *vzhledem k typu statistických vlastností*. Statistické vlastnosti mohou být reprezentovány v různém stupni:

(1) Projekt byl mezinárodní, zúčastnilo se ho patnáct zemí. Přestože metodologie byla předem podřízena jednoznačně určenému komparativnímu cíli, byly i v rámci projektu číněny pokusy o zobecnění na uvedené základní soubory. Tyto závěry pochopitelně nemusí být platné, neboť jejich validita není prokázána.

(2) Pro takové typy dotazů jsou vyvíjeny různé typy otázek, jejichž cílem je uvedené nebezpečí odstranit. Jde především o různé modely takzvaných *randomizovaných otázek*. U nás tyto postupy nebyly zatím dostatečně podrobně vyzkoušeny ani nebyly publikovány pokusy o jejich modifikace, přizpůsobení našim podmínkám či rozvíjení.

- a) procentní zastoupení;
- b) průměry a rozptyly;
- c) jednorozměrné rozložení četností a jeho tvar;
- d) vzájemné vztahy mezi dvojicemi proměnných měřené statistickými koeficienty;
- e) dvojrozměrná statistická rozložení (kontingenční tabulky);
- f) korelační matice a výsledky jejich zpracování (diskriminační analýza, faktorová analýza, regresní analýza apod.);
- g) mnohorozměrná rozložení četností;
- h) celá statistická struktura všech proměnných (mnohorozměrné rozložení všech znaků v dotazníku).

Rozhodnutí o reprezentativitě vzhledem k těmto vlastnostem je velice obtížné. Požadavek jejího zajištění úzce souvisí s úlohou, jejím obsahem a cílem, pro nějž data sbíráme a analyzujeme. Zajištění reprezentativity jednoho typu vlastností vůbec nelze vztáhnout na jiné typy. Tak například reprezentativnost průměrů neznamená reprezentativnost celého rozložení, správná procenta a jednorozměrná rozložení neznamenají, že odpovídající měření vztahů mezi proměnnými jsou také správná. A naopak: I když jednorozměrná rozložení jsou zkreslená, může odhad koeficientů asociace či regresní křivky být správný.

V plánech a projektech sociologických šetření jsou tedy velké rezervy efektivního uvažování. Dobře formulovaný cíl a přesné zaměření úlohy a požadavků k tomuto cíli mohou podstatně snížit náklady, neboť nemusíme v terénu zajišťovat více, než je prakticky vyžadováno. Proto se liší výzkumy veřejného mínění, jejichž úkolem je zjistit stav názorů a mínění určité populace v určitém období, od složitých komparativních výzkumů, jejichž cílem jsou explanativní závěry a postižení podstaty sociálních jevů.

Nejlépejším chápáním reprezentativnosti je *reprezentativnost vzhledem ke stupni pokrytí cílové populace*. Můžeme totiž reprezentovat populaci úplně, můžeme vynechat její nepočtenou a špatně dostupnou část (většinou pro statistické závěry nepodstatnou) či některé závažné skupiny. *Nepokrytí části populace může být:*

- záměrné,
- nechtěné, ale ne nebezpečné vzhledem k malému efektu,
- nechtěné a identifikované, a tudíž interpretačně nebo statisticky korigovatelné,
- neidentifikované, neznámé, a tudíž interpretačně nebezpečné.

K záměrnému vynechání patří typické případy vypouštění malých, špatně dostupných skupin obyvatelstva, skupin, které jsou mobilní, či skupin, jejichž dosažení je příliš nákladné. Vždy je však nutno zvážit, jaký význam má vypuštění skupiny pro závěry. Málo početná skupina nemůže příliš změnit výsledné výběrové průměry ani procenta, pokud se ovšem podstatně neliší od ostatních. Tak při výzkumech průmyslové sociologie (při zkoumání vlivu techniky) se ukázalo, že v mnoha dílnách existuje nepočtená skupina vysoce kvalifikovaných dělníků, jejichž názor se však u některých otázek diametrálně liší od názoru ostatních především proto, že je hlubší, kvalifikovanější a promyšlenější. Vypuštění těchto několika málo respondentů, kteří tvoří nepatrné procento celé skupiny, může vést ke značné ztrátě informace a k podstatnému zkreslení výsledků. Jiným příkladem je nezahrnutí některých etnických skupin z důvodů jazykových.

Největším nebezpečím ovšem je působení *samovýběru*. Lidé, kteří mají určitou vlastnost, odmítají odpovídat (například lidé se záporným postojem k určité věci nebo k tazateli raději odmítnou odpovídat na dotazník, než aby svůj postoj vyjádřili). Tak se též projevuje mechanismus ankety, která je ve většině případů též nezobecní-

telná. Na anketu odpovídají lidé, kteří očekávají odměnu, či lidé, kteří rádi píší nebo rádi dávají připomínky apod. Toto nebezpečí se projevuje i u kvótních výběrů, a to ve formě nekontrolovaného a nekontrolovatelného samovýběru respondentů či zaměřeného výběru respondentů tazatelem. (Tazatel si vybírá určitý typ respondentů, především takové osoby, které jsou ochotny odpovídat a s nimiž snadno naváže kontakt.) Tím vzniká neidentifikované zkreslení reprezentativity.

Významné aspekty reprezentativity tedy jsou:

- obsahová reprezentativita úlohy a instrumentu;
- reprezentativnost vzhledem k cíli a úkolu;
- reprezentativnost vzhledem k typu statistických vlastností;
- reprezentativnost vzhledem k pokrytí cílové populace.

Posledním aspektem se budeme dále zabývat podrobněji, neboť v tomto smyslu se obvykle o reprezentativitě mluví a takto je ve výzkumné praxi většinou pojem intuitivně chápán. Budeme tedy dále předpokládat, že jednotlivé znaky a výzkumná instrumentace jsou jednoznačně dány (dotazník, záznamový list apod.), že abstrahujeme od chyb měření, přenosu a zpracování dat. Půjde tedy o reprezentativitu ve smyslu vztahu výběrového souboru a cílové populace.

2. Výběrový postup a výběrový soubor

Výběrový soubor vzniká jako realizace výběrového postupu (výběrového plánu); má své výběrové uspořádání. Výběrový postup může být velice jednoduchý, či naopak velice komplikovaný, komplexní. Výběrová uspořádání odpovídají výběru buď pravděpodobnostnímu, kvótnímu, nebo záměrnému. Ty jsou značně rozdílné, což je dáno principem jejich skladby i typem reprezentativnosti, kterou představují. Pravděpodobnostní výběry byly prací společenskovědních oborů i ostatních věd prověřeny jako nejsprávnější a jako základní pro zajištění reprezentativity. Skládají se z různých postupů, respektive technik:

- a) postup prostého náhodného vybírání;
- b) postup systematického vybírání;
- c) vybírání skupinek;
- d) stratifikace;
- e) vybírání ve stupních;
- f) fázové vybírání;
- g) vážení při vybírání jednotek.

Různé modifikace pravděpodobnostního výběrového postupu a odchylky od něj jsou zaváděny buď z důvodů zvýšení přesnosti a reprezentativity (například řízený výběr), či z důvodů technických (například "snow-ball" výběrový postup), respektive z ekonomických (například kvótní výběr). Za určitých velmi striktních předpokladů lze na uvedené modifikace aplikovat výsledky modelů pravděpodobnostních výběrů, přestože tyto typy pravděpodobnostní nejsou (v některých případech se však jako pravděpodobnostní chovají).

Jednotlivé výše uvedené složky mohou na sebe různě navazovat a různě se prolínat. Každý typ uspořádání má pak svůj vlastní způsob odhadu populačních parametrů a jejich směrodatných chyb. Časté chybné závěry jsou způsobeny právě tím, že i při komplexním výběrovém uspořádání používáme nevhodné odhadové vzorce (většinou jsou to vzorce pro nejjednodušší prostý náhodný výběr).

Naše úvahy musíme vztáhnout buď k výběrovému plánu, k výběrovému uspořádání, tj. ke způsobům, metodě, postupu, nebo k realizaci metody, k naplnění plánu, ke konkrétnímu souboru jednotek, jenž je nositelem výzkumné informace, tj. k vý-

běrovému souboru. Tyto dva různé pohledy, které jsou oba velmi důležité, nutno na jedné straně rozlišovat, ale zároveň je nutno vždy je hodnotit současně.

I nekorektní způsob a plán může (i když málo pravděpodobně) dát reprezentativní soubor. Naproti tomu i ten neoptimálnější, nejvhodnější, nejlépe organizovaný výběrový postup může dát zkreslující výsledný výběrový soubor (to však je málo pravděpodobné).

Výběrové uspořádání je plánováno především z hlediska zajištění reprezentativnosti a přesnosti odhadů v takovém stupni, který vyžaduje úloha, a zároveň z hlediska praktické uskutečnitelnosti a ekonomičnosti. Do rozhodování o výběrovém uspořádání vstupují tedy dva často konfliktní aspekty: *zajištění co nejnižších nákladů a co nejvyšší přesnosti*. Uspořádání musí odpovídat technickým možnostem, aby nedocházelo k nekontrolovaným hrubým porušením postupu, jež by mohly pokazit záměry a porušit reprezentativnost.

Komplexní výběrový postup má v praxi obvykle velice složitou teoretickou strukturu, která přispívá ke zvýšení přesnosti výsledků a zjednodušení postupu při praktickém vybírání. Teoretická složitost úlohy a jejího řešení tedy nevede v praxi ke komplikacím. Vyžaduje však, aby pro vyhodnocení souboru dat byly používány speciální formule jak pro odhad průměrů a procent, tak pro odhad jejich směrodatných chyb a konfidenčních intervalů. To ovšem není tak velký problém, je-li k dispozici počítač. Postupu však musí být věnována dostatečná pozornost, která v praktické výzkumné práci často chybí.

V komplexních výběrových postupech se objevují především různé aplikované postupy stratifikace, vybírání skupinek a vybírání ve stupních. Aplikujeme často systematický výběr, abychom zajistili stratifikační efekt, vybíráme proporcionálně k daným velikostem apod. Mnohdy chceme uspořádat výběr tak, aby sloužil dvěma cílům: reprezentoval jedince i jejich určité skupinky.

Příkladem může být výběr školních dětí v ČSR. V každém kraji vybíráme zvlášť (oblasti), a to proporcionálně počtu školních dětí v krajích. V Praze vybíráme ze seznamu všech škol, v každém kraji pak vybíráme nejprve tři okresy, v každém okrese dále školy a v každé škole jednu nebo více tříd. Vybranou třídu můžeme zahrnout do výběru celou nebo její část — přitom vybíráme buď skupinku o dané velikosti, nebo danou část třídy. Na každém kroku musíme přesně určit pravidla výběru — buď volíme pravděpodobnosti stejné, či jsou pravděpodobnosti proporcionální velikosti skupinek. Obdobně můžeme provést výběr v SSR. Vzhledem ke komparativním cílům pak mohou být v ČSR a SSR výběry stejně velké. Mohli bychom uvést další varianty tohoto příkladu — každá z nich však pro praxi musí být účelná z hlediska přesnosti výsledků, praktičnosti nebo ekonomičnosti.

Výběrové šetření můžeme vyhodnotit jenom tehdy, jestliže vytvoříme *matematický model postupu*, který určuje všechny složky výběrového uspořádání a umožňuje určení výsledných výpočetních formulí a výpočet pravděpodobností pro výběr každé jednotky na všech stupních. Matematický model odráží postup kontroly náhodnosti tak, jak ji zavádíme ve výběrovém uspořádání.

Základní otázkou ovšem zůstává: Jakou populaci chceme reprezentovat? V příkladě výběru z populace školních dětí musíme určit, zda půjde o populaci celou, či věkově omezenou, zda se omezíme na populaci školních dětí z určitého typu škol, či zda nás zajímá základní soubor tříd, pionýrských organizací, zájmových kroužků, či dokonce základní soubor škol jakožto samostatných organizačních jednotek. Takové statistické skupinky nás zajímají jako sociologické celky, které mají určitou organizaci, sociální roli, jsou sociologicky či jinak výzkumně důležité. A konečně nás může zajímat obojí: přehled o dětech i fungování života ve třídách; nebo dokonce několik takových souborů a navíc například i populace učitelů.

Přitom vzniká řada intuitivních otázek:

- Je soubor reprezentativní pro ČSR a pro SSR zvlášť?
- Je reprezentativní pro ČSSR, jestliže jsme volili stejné výběrové velikosti pro obě republiky?
- Neměli bychom vynechat část respondentů, abychom zajistili proporcionální zastoupení obou republik (není však škoda nevyužít nákladně shromážděných informací)?
- Máme reprezentativně pokryté jednotlivé kraje?
- Co dělat, když náhodný postup nezahrnul určitý okres, o němž víme, že je tam zvláštní, netypická, velmi odlišná situace?
- Můžeme každou školu, či dokonce třídu reprezentovat jakožto dílčí populaci?
- Musí být zajištěna reprezentativita každého stupně, aby byl celý výběrový soubor reprezentativní?

Je zřejmé, že jednoduchý průřez tříděním i tabelacemi není pro ČSSR reprezentativní, pokud nebudou data o ČSR a o SSR zcela statisticky homogenní; v celém souboru se projeví zvýšená váha na výsledky SSR. Tabelace nemusí být reprezentativní ani pro žádnou z později definovaných částí populace. Má tedy smysl dělat takové složité plány? Co dělat, jsme-li do takových postupů (zdánlivě nereprezentativních) prostě přinuceni objektivními okolnostmi a podmínkami správného uspořádání, dostupností výběrových opor, komplexem sociologických záměrů a otázek?

Smysl teorie pravděpodobnostního výběru je v tom, že buduje modely a matematicko-statistickou teorii různých výběrových uspořádání, která *umožňují pro daný postup určit korespondenci mezi výběrovým souborem a cílovou populací*. Tyto modely jsou tedy základem pro zobecňování, pro zjištění stupně neurčitosti, který při zobecňování musíme respektovat. Odvozeně z nich plynou adekvátní odhadové vzorce, na nichž zobecňování z výběrového souboru na cílovou populaci závisí. Čím komplexnější je uspořádání, tím složitější jsou tyto vzorce. Přímé vypočítané průměry a přímé tabelace nemusí vypovídat o situaci v základním souboru. U komplexních přístupů je většinou nutné používat různých vah, respektive korekcí.

Pojem reprezentativity je často zjednodušeně chápán právě jako přímá možnost použití výběrových (nevážených) procent a průměrů jako charakteristik základního souboru. Existují sice postupy, které tento požadavek splňují, no vždy jsou však výhodné a ekonomické. Celkově je ovšem takové zjednodušené chápání reprezentativity chybné. Musíme tedy zhodnotit dvě hlediska: *reprezentativitu výběrového uspořádání a reprezentativitu výběrového souboru* (jakožto realizace výběrového plánu). Oba aspekty určují stupeň zobecnitelnosti dat a metody jak toto zobecnění provádět (odhadové vzorce). Současně vytvářejí *reprezentativitu výběrového šetření*.

V dalším budeme vždy předpokládat, že reprezentativnost dat váže oba uvedené aspekty. Reprezentativnost výběrového postupu je teoretická vlastnost modelu šetření, reprezentativnost souboru je vlastnost naplnění modelu praktickou realizací.

V případech, kdy máme k dispozici záznam původních dat (na magnetické páse, děrných štítech, respektive v dotaznících apod.), potřebujeme nutně popis modelu výběrového postupu, popis metody sběru, tj. dokumentaci, která dobře charakterizuje genezi daného datového souboru. Data sama vypovídají velice málo, pokud nevíme, jak je vhodně zpracovat, abychom dostali zobecnitelné tabelace a adekvátní odhady. V této souvislosti je patrný význam metodologické dokumentace každého výzkumu, který má být archivován, distribuován, sekundárně použit, či má sloužit jako charakteristika jednoho časového průřezu v dlouhodobém sledování daných jevů.

Proto je také důležité uvést metodologické aspekty výzkumu ve výzkumné zprávě a při oponentuře, aby bylo prokázáno, že výsledné charakteristiky dat byly správné

spočítány a že v analýze byly použity vhodné metody odhadu, zajišťující reprezentativitu výsledků.

Přímá reprezentativnost, tj. přímý přenos výsledků z datového souboru na cílovou populaci, je vázána na situace, kde výběrové charakteristiky jsou *nevychýlenými* (nestrannými) *odhady* svých populačních protějšků nebo (u velkých souborů) konzistentními odhady. O těchto pojmech zde nepojednáme, čtenáře odkazujeme na citovanou literaturu. Taková přímá zobecnitelnost je možná jen tehdy, když každá jednotka populace má stejnou pravděpodobnost, že se do výběru dostane. Tuto vlastnost má celá řada i velmi komplexních uspořádání. Jinak vypadá situace, chceme-li odhadnout výběrové chyby pomocí elementárních vzorců, využívajících celkové standardní odchylky spočítané přes data výběru. To je možné udělat pouze pro prostý náhodný výběr, respektive pro některý jeho ekvivalent. Taktéž běžné statistické testy, jako je χ^2 -kvadrát a další, budou mít platnost jen pro prostý náhodný výběr a pro případy, kdy jedním ze znaků tabulky je znak stratifikační. U komplexních postupů musíme vždy počítat standardní chyby odhadu závislé na uspořádání, tj. odvozené z modelu.

U metod odhadu můžeme volit různé přístupy, které nám pomáhají informaci upřesnit a korigovat některá zkreslení. Jde o přístupy používající dodatečné informace. Jsou to především metody poměrového s regresního odhadu; jejich použití vede ovšem k jiným standardním chybám odhadů.

Zúžené chápání reprezentativnosti souboru, které znamená pouze přímý přenos tabelací a výsledků z datového souboru na celou cílovou populaci, je tedy zřejmě zploštěním pohledu na statistickou práci, statistickou metodu a zploštěním pohledu na celý metodologický problém. Jde o chápání laické, které by nutně mělo z naší sociologie vymizet. Vede k nevyužití možností statistické metodologie i ke zbytečně vysokým nákladům, popřípadě i ke zkreslení, respektive k systematickému zkreslování výsledků, které však není známo a odhaleno.

Nakonec uvedme ještě jeden aspekt reprezentativnosti. Reprezentativnost pravděpodobnostních výběrů má svůj *neurčitostní charakter*: kterýkoli ukazatel u kteréhokoli znaku se nemusí přesně krýt s populačními charakteristikami. Nikdy nemůžeme očekávat jednoznačnou shodu mezi výběrem a cílovou populací. Jestliže taková shoda opravdu nastane, budeme spíše překvapeni. V čem tedy spočívá reprezentativita? Je to především aspekt řízené náhodnosti, který ji zajišťuje, a vybíráme-li například jednoprocentní soubor prostým náhodným výběrem, pak předpokládáme, že vybraný soubor bude v rámci předepsaných tolerancí zhruba takový, aby umožnil dostatečně přesné odhady populačních charakteristik (pravděpodobnost velkých odchylek je malá). (3) Takováto pravděpodobnost „postačující shody výběru s po-

(3) I v neoptimálnějším pravděpodobnostním výběru funguje náhoda, i když je to náhoda řízená. Její předností je výběr objektivizovaný, jejím nedostatkem je, že ve svých konečných fázích ji nemůžeme kontrolovat. Když budeme podle tabulky náhodných čísel vybírat ze souboru 25 lidí výběr o pět osobách, pak každá pětice má stejnou šanci, že se do výběrového souboru dostane. Máme-li například základní soubor složený z pěti žen a dvaceti mužů, pak je možné, že dostaneme výběr obsahující právě oněch pět žen, tedy výběr naprosto nereprezentativní. Pravděpodobnost tohoto jevu existuje, je však nesmírně malá. Největší šanci na výběr má takový soubor, u něhož bude poměr mužů a žen 1 : 4 nebo relace málo odlišné. Takových souborů je totiž daleko více, a proto mají největší

šanci, bohužel však ani výše uvedená extrémní odchylka nemá šanci nulovou.

Pro poměr žen 5 : 0 je pravděpodobnost výběru $P = .00001882$. Pro poměr 1 : 4 (odpovídá výše popsané populaci) dostaneme pravděpodobnost $P = .455957$. To je způsobeno tím, že prvnímu poměru odpovídá pouze jeden výběrový soubor z 53 130 možných výběrových souborů (o velikosti 5), zatímco druhému poměru odpovídá 24 225 pětic, v nichž je právě jedna žena a čtyři muži. Tyto poměry se ještě více vyostřují při velkých souborech. Vždy ovšem zůstává nutná fluktuace kolem skutečné hodnoty, tj. vysoká možnost realizovat soubor s poměrem blízkým skutečné hodnotě. Právě tuto fluktuaci charakterizují konfidenční intervaly.

populací“ poroste s velikostí výběru,(4) s vhodnou stratifikací, s optimálním rozmístěním celkové výběrové velikosti do oblastí a s vhodným poměrem velikostí skupinek a jejich vybraných částí. Bude klesat s použitím výběru homogenních skupinek, s počtem stupňů a se zvyšující se heterogenitou vlastní cílové populace. Reprezentativnost se bude dále snižovat s počtem zkoumaných proměnných, neboť se zvyšuje nebezpečí, že u některých z proměnných se projeví extrémní a interpretačně nebezpečný odklon.

Reprezentativita souboru, dat, postupu, šetření tedy zřejmě není (a nemůže být) dichotomický pojem ve smyslu „reprezentativní — nereprezentativní“. Musíme mluvit o stupni reprezentativnosti, o stupni, v jakém můžeme data zobecňovat na předem danou cílovou populaci.

3. Vymezení pojmu reprezentativita

Shrneme-li předcházející diskusi pojmu, vidíme, že reprezentativita či použitelnost dat může být ohrožena v celé řadě kroků, jimiž jsou i formulace záměru a cíle (tílohy, tématu, populace) — určení vlastností — operacionalizace vlastností, konstrukce výzkumných instrumentů — plán výběrového uspořádání — realizace výběrového plánu — sběr informace — určení výběrového souboru — zpracování dat — interpretace.

Omezíme-li se pouze na vztahy mezi cílovou populací a výběrem, schéma se zjednoduší: cílová populace — plán výběrového uspořádání — výběrový soubor — metoda zpracování dat. V těchto krocích odhlížíme od chyb měření, chyb vnesených tazateli, kódováním, přenosem dat atp. I tak jde o velmi náročný a nesmírně komplikovaný postup. Není náhodou, že se ve statistice vyvinula samostatná specializace, která se zabývá teorií a praxí výběru a všemi aspekty výběrových šetření. Zatímco výzkumník hodnotí reprezentativitu v podstatě intuitivně, specialista-statistik ji musí posoudit ze všech uvedených hledisek i z mnoha dalších.(5)

Náš pojem můžeme tedy charakterizovat takto:

Reprezentativnost (reprezentativita) výběrového šetření je stupeň přenositelnosti informace obsažené ve výběrovém souboru na cílovou populaci či její zvolenou část; tato zobecnitelnost je zprostředkována statistickým modelem výběrového uspořádání a znalostí odchylky předpokládaného modelu od její realizace. Zahrnuje stupeň obsahového pokrytí úlohy (počet a role těch znaků, jejichž informace je zobecnitelná, stupeň pokrytí statistických vlastností dat (vlastností statistických měr, distribucí a celkové struktury dat) a rozsah pokrytí populace (zahrnutí všech typů jednotek populace do výběru). Je to stupeň zachycení populační variability rozložení všech proměnných výzkumného instrumentu pomocí daného realizovaného výběrového souboru a daného modelu výběrového uspořádání.

(4) Přesnost odhadu roste s odmocninou velikosti výběru, a to v tom smyslu, že jestliže se velikost zvětší k -krát, konfidenční interval se zmenší $\sqrt{k}$ -krát. Například dvojnásobné zvýšení výběru znamená 1,41násobné zmenšení konfidenčních intervalů, čtyřnásobné vyšší výběrová velikost znamená poloviční konfidenční intervaly.

(5) Existuje chybný názor, že teorii a praxi výběru se statistik naučí po přečtení jedné či dvou monografií a že jde v podstatě o rutinní práci. To je pravda pouze z malé části. Metody hodnocení reprezentativity jsou jen v malé míře jednoznačně předepsané a snadno

aplikovatelné, každý výzkum vyžaduje velmi citlivou reflexi ve výběrovém plánu. Profese statistika specializujícího se na výběrová šetření vyžaduje velké teoretické znalosti a praktický cit. Profese většinou vyžaduje i dlouhou praxi. Mnohdy se postupy zdají velice jednoduché a vypadají jako opakování jiných výzkumů — přesto je práce na nich a cesta k takovému jednoduchému výsledku značně komplikovaná. Statistik musí modelově připravit a prověřit možnost co nejjednoduššího řešení. A to nepatří vždy k nejjednodušším výzkumným rozhodnutím.

Toto obecné vymezení charakterizuje pojem co do jeho obsahu, vymezuje jeho roli ve výzkumném procesu i jeho užitečnost. Je zřejmé, že pojem reprezentativity je jedním ze základních metodologických pojmů, charakterizujících kvalitu našich závěrů. Připomeňme, že zde mluvíme především o statistických zobecněních — **ná-
vazný** interpretační krok však je s výsledky statistického zobecnování těsně spojen.

Dalším nutným metodologickým úkolem je operacionalizace pojmu, tj. nalezení zpracovatelných charakteristik, které reprezentativitu určují a umožňují jednoduchou a efektivní práci s tímto pojmem. Současně je třeba určit faktory, které reprezentativitu snižují a narušují v praktickém procesu výzkumu, a nalézt korektivy, které umožní některé nepříznivé působící vlivy eliminovat.

4. Dva základní ukazatelé odklonu od reprezentativity: neurčitost odhadu a vychýlení

Reprezentativnost můžeme operacionalizovat pomocí dvou základních ukazatelů: *neurčitost odhadu a vychýlení odhadu*. Tyto dva ukazatelé umožní měřit stupeň reprezentativnosti výběrových souborů dat. Oba charakterizují přesnost statistických výsledků, které výběrovým postupem dostáváme; jedná se o přesnost vzhledem k cílové populaci.

Označíme si:

A = skutečná hodnota parametru v populaci (například průměr, rozptyl, procento, korelační koeficient, koeficient asociace, průměr poměrů, poměr průměrů, regresní koeficienty atd.),

a = odhad parametru A (tj. hodnota, kterou dostáváme pomocí příslušných vzorců z výběrových dat a která reprezentuje neznámou hodnotu A),

s_a = směrodatná chyba odhadu a (hodnota získaná z příslušných vzorců vyplývajících z modelu postupu),

a^* = očekávaná hodnota odhadu a (tj. průměrná hodnota odhadu a ze všech možných realizací výběrového souboru, které odpovídají danému zvolenému modelu).

Směrodatná chyba odhadu a je mírou neurčitosti tohoto odhadu, mírou jeho nejednoznačnosti, jeho rozptýlenosti kolem očekávané hodnoty odhadu a^* . Čím větší je s_a , tím méně přesné je a . Odhad a tedy charakterizuje neznámou hodnotu a^* . Jestliže $a^* = A$, říkáme, že a je nevychýlený odhad (respektive nestranný odhad) populační charakteristiky A .

Hodnota $B = A - a^*$ se nazývá vychýlení odhadu. V případě, že $B \neq 0$, dostáváme systematicky vyšší nebo systematicky nižší hodnoty pro parametr A . Toto vychýlení se nemusí projevit v jednom šetření.⁽⁶⁾ V dlouhodobém průměru (při opakování) se vychýlení B projeví jasně.

O hodnotě a^* vypovídáme zatím pouze s neurčitostí, která vyplývá z s_a (tj. z modelu i naplnění výběrového postupu). Abychom věděli, s jakou přesností můžeme zobecňovat, vytváříme tzv. *intervaly spolehlivosti (konfidenční intervaly)*. Tyto intervaly odpovídají předem zvolené spolehlivosti, s níž chceme zobecňovat. Je přirozené, že čím větší spolehlivost zobecnění budeme vyžadovat, tím větší interval pro skutečnou hodnotu a^* dostaneme. Tvar intervalu je

$$a \pm t \cdot s_a$$

t je kvantil Studentova rozložení s ($n - 1$) stupni volnosti, kde n je počet pozorování

(6) „Systematicky vyšší“ nebo „systematicky nižší“ je zde chápáno ve smyslu střední hodnoty. Ve skutečném výběru může hodnota padnout zcela na opačnou stranu od A , než indi-

kuje vychýlení B . Při opakování v šetřeních však bychom dostali hodnoty s průměrem blízko a^* , nikoli blízko A .

(údajů) vstupujících do odhadu. Tento údaj snadno nalezneme(7) v běžných statistických tabulkách. Proto je možnou mírou neurčitosti délka konfidenčního intervalu ($= 2 ts_a$). Chceme-li abstrahovat od proměnlivého t , které je závislé na volitelném procentu spolehlivosti, pak s_a je vhodná míra neurčitosti odhadu.

Obě hodnoty s_a , B charakterizují tedy odklon od reprezentativity. Oba aspekty se snažíme eliminovat.

1. Neurčitost odhadu redukuje hlavně tím, že volíme vhodné uspořádání výběru a jeho dostatečnou velikost. Zároveň i tím, že zvolíme vhodné metody odhadu. Tyto tři faktory mají na neurčitost vliv, mohou ji zvětšit, ale především — při vhodném použití — redukovat. Jak už bylo uvedeno, s_a klesá s odmocninou velikosti výběru. Největší rezervy pro zpřesnění jsou ve vhodné stratifikaci.
2. Vychýlení je o to nebezpečnější, že jej obtížně identifikujeme.(8) Vzniká především nepokrytím celé populace, nerespektováním výběrového uspořádání při volbě odhadových vzorců a dále i volbou metody odhadu (odhady regresní a poměrové).

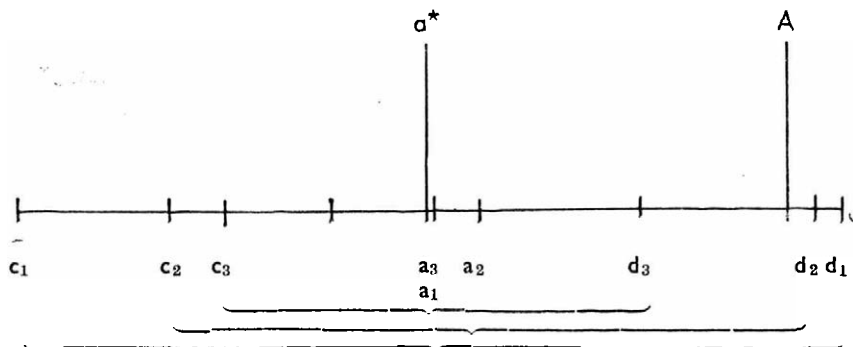
Vychýlení můžeme redukovat, respektive eliminovat, především dodržením zvoleného a teoreticky zdůvodněného uspořádání výběru a správnou volbou odhadových vzorců. Počet pozorování má vliv pouze na vychýlení způsobené metodou odhadu (u odhadů regresních a poměrných).

Vychýlení způsobené nepokrytím populace (například odmítáním respondentů, které má systematickou příčinu) a vychýlení způsobené neadekvátní volbou odhadových vzorců (například nerespektování různých vah, s nimiž jsou jednotky vybírány) nelze odstranit zvětšením velikosti výběrového souboru!

Naopak tím, že se neustále snižuje neurčitost odhadu, a se s vyšším počtem pozorování blíží k hodnotě a^* , konfidenční interval se zužuje a tím spíše hodnotu A (pokud je tedy rozdílná od a^*) nepokryje (B zůstává konstantní).

Situace je znázorněna na obrázku 1: Nejprve získáme průměr a_1 a jeho konfidenční interval (c_1 , d_1) z malého výběru. Interval pokryje a^* i A . Při zvětšení souboru dostaneme a_2 a jeho interval (c_2 , d_2), obě hodnoty a^* , A jsou pokryty. Při dalším zvětšení výběru je a_3 přesnější, interval (c_3 , d_3) užší, A už v něm neleží.

Obrázek 1



(7) Pro větší počet pozorování, řekněme pro více než 150, je toto číslo velmi blízké normálnímu kvantilu. V praxi se používá většinou spolehlivost 95 %, skóre normálního rozložení je 1,96; pro rychlé výpočty bereme většinou $t = 2$, což vyhovuje i výběrům které obsahují více než 50 pozorování. Pro větší n je ovšem

výsledný interval nepatrně menší. Proto 95 % intervaly spolehlivosti se většinou uvádějí jako $a \pm 2s_a$, přičemž šířka konfidenčního intervalu je $4s_a$.

(8) V mnoha případech se pokoušíme odhalit alespoň horní hranici jeho absolutní hodnoty či jeho směr.

Příklad:

Chceme zjistit % lidí, kteří mají určitý názor. Předpokládejme, že tento názor má 70 % lidí (toto číslo ovšem neznáme). Předpokládejme, že lidé odmítají odpovídat podle mechanismu: odpoví 90 % z těch, kteří mají daný názor a 20 % z těch, kteří mají názor opačný. Celkem se tedy populace zredukuje na $(.7 \times .9 + .3 \times .2) \cdot 100 \% = 69 \%$. V této části populace je však poměr kladných odpovědí k záporným 91 % : 9 % (oproti původnímu procentnímu poměru 70 % : 30 %). Zde tedy $A = 70 \%$, $a^* = 91 \%$, $B = 21 \%$ (!). Ve výzkumné situaci tato čísla neznáme a jsme pochopitelně odkázáni pouze na empirická měření. Odhadujeme však hodnotu 91 %.

Tak ve výběru o velikosti $n = 100$ můžeme například zjistit $a = 85 \%$. Směrodatná chyba pro tento odhad je 3,6 %, a tedy pětadevadesátiprocentní interval spolehlivosti je $85 \% \pm 7,2 \%$, tedy zhruba (78 %, 92 %). Jak je vidět, hodnota a^* (kterou ovšem neznáme) do intervalu padne, hodnota A ovšem nikoli.

Kdybychom zvýšili počet na $n = 1000$ a dostali například $a = 90 \%$ (očekáváme větší přesnost, a proto a by mělo být blíže k a^* než pro $n = 100$), dostaneme $s = 0,95$, a tedy 95% interval spolehlivosti jako $90 \% \pm 1,9 \%$, což je přibližně (88 %, 92 %). Tento výsledek pokrývá daleko těsněji hodnotu a^* , je však vzdálenější od skutečné hodnoty A .

Příklad odpovídá běžným zkušenostem ze šetření, v nichž respondenti neradi projevují určitý názor či postoj. Je reálný (i když neočekávaně nepříjemný).

Zvyšování velikosti souboru tedy nevede k odstranění vychýlení, naopak vede k přesnějšímu měření chybné hodnoty a^* .

Metody odstranění či redukce vychýlení závisí na jeho příčinách; nejlepší cesta spočívá v preventivním vyloučení vychylovacího mechanismu, je-li to možné. Opět zde nebudeme roznědět obsahové vychýlení (nevalidní otázky, špatně volené proměnné, nevhodné operacionalizace).

Příčiny vychýlení tkví v nedostatečné metodologické přípravě instrumentu, v nedocení tohoto typu chyby, v nepokrytí populace plánem šetření, v odmítnutích a ve fungování samovýběru, respektive zkreslujícího vlivu výběru respondenta tazatelem (to je hlavní argument proti kvótním výběrům v té formě, jak se aplikují u nás).

Typickým příkladem nedostatečného pokrytí cílové populace je výběr předplatitelů ve výzkumech čtenářských populací, chceme-li charakterizovat celou čtenářskou obec. Bylo už empiricky prokázáno, že předplatitelé nemohou čtenářskou obec reprezentovat (zejména není-li to pro každý jednotlivý časopis zvlášť vyzkoušeno). U odhadu poslechovosti rozhlasu, televize či odebrání časopisu nemůžeme zobecňovat soubor dat získaných anketou či z dobrovolníků, byť by z anket byly vybrány osoby tak, aby splňovaly nejružnější kvóty. Zde působí samovýběr založený na motivaci a postojích, což vede k nadhodnocení procenta poslechu, neboť u lidí, kteří více poslouchají a mají lepší poměr k rozhlasu, je tu vyšší pravděpodobnost, že na anketu odpoví a že se o ní vůbec dovědí. Právě v této oblasti je otevřené pole pro metodologický výzkum a metodologické zpracování dat, především pro ty instituce, které pracují s panelovými soubory či mají vlastní tazatelské sítě, a pro ty situace, v nichž vznikají soubory samovýběrem či jsou jím silně ovlivněny (dobrovolníci do panelu, zájemci, placení či jinak motivovaní respondenti).

Typický je též příklad, v němž je nutno data vážit předem určenými vahami, aby-
chom dostali nevychýlené odhady. Vybíráme-li respondenty při projevu v nějaké aktivitě, například při návštěvě koncertů, výběrový soubor návštěv bude dobře reprezentovat „průměrný koncert“. Zajímá-li nás však cílová populace návštěvníků

koncertu, pak musíme korigovat výsledky. Lidé, kteří danou aktivitu provozují častěji, mají větší šanci dostat se do výběru. A jelikož intenzita provozování aktivity je jistě korelována s nejrůznějšími znalostmi, názory, postoji a přáními, při každém použití souboru návštěv pro zobecnění na obec posluchačů bychom dostali silně zkreslené výsledky, a tudíž i zcela nesprávné interpretace a popřípadě i chybná následná důležitá rozhodnutí. Tento fakt byl už empiricky prokázán. Chybu lze eliminovat vhodnou přípravou instrumentů a zařazením proměnných, které by indikovaly, respektive přesně určily váhy pro převod. Je však nutné počítat s tímto aspektem předem a sestavit vhodný model výběru, z něhož plynou i postupy zpracování.

Vychýlení je tedy velmi nepříjemné porušení reprezentativity. Jsou však situace, kdy vychýlení nevadí. Je to v případě, že funguje stejně na různých částech populace a nás zajímá porovnání dvou či více částí. V rozdílu průměru se totiž vychýlení vyruší:

$$a_1^* - a_2^* = (A_1 - B) - (A_2 - B) = A_1 - A_2$$

a odhad rozdílu: $a_1 - a_2$ je vhodným odhadem pro rozdíl $A_1 - A_2$ (indexy odpovídají první a druhé části souboru dat, respektive první a druhé části zkoumané populace).

Chceme-li dostat srovnatelné údaje (například v panelu či při opakovaném výzkumu pro zjištění časového vývoje) a první část šetření byla chybně koncipována s vychýlenými výsledky, pak jestliže velikost vychýlení není známa, musíme se při opakování dopustit stejné chyby, aby se nám vychýlení při komparaci výsledků zrušilo.

Při sekundárních analýzách je nutno znát okolnosti výběru a výběrové plány, aby nebyly interpretovány metodologické chyby a artefakty jako závažné výsledky.

5. Výběru a jejich reprezentativita

Zatím jsme provedli obecnou diskusi, která sice vycházela z pravděpodobnostních výběrů, platila však obecně. I když je metoda pravděpodobnostního výběru základní pro práci sociologického výzkumu i pro výzkumné účely jiných věd, používáme další dva časté typy: záměrný výběr a kvótní výběr. Optimální a komplexní pravděpodobnostní výběr není v praxi vždy možno aplikovat. Výzkumná situace nás nutí k různým modifikacím a kompromisům. V rozhodování se projevují především požadavky reprezentativnosti, náklady, časová omezení, praktické možnosti a organizační zabezpečení. Konkrétní rozhodnutí musí brát v úvahu všechny tyto okolnosti.

Jsou případy, kdy je jedno hledisko jednoznačně dané: požadavek vysoké reprezentativnosti a důležitosti přesnosti výsledků může vést (bez ohledu na cenu, nutnost dostatečného časového předstihu a organizační úsilí) k vytváření speciálních seznamů pro dané šetření. Jsme-li omezeni v nákladech, musíme volit levnější uspořádání na úkor přesnosti a reprezentativnosti. Nemáme-li výzkum zabezpečen organizačně, musíme se obrátit na některou z existujících tazatelských sítí a popřípadě se přizpůsobit rutinní metodice sběru dat v této síti, která nemusí být pro naše účely vhodná.

Praktických potíží je mnoho. Malý důraz se například klade na kvalitu informace vzhledem k cíli, k němuž má sloužit. Často se opomíjí apriorní úvahy o nutnosti kvalitní informace, příprava snadno a rychle končí rozhodnutím o organizačně nejspokojnější cestě, — aniž by se hledala cesta sice méně schůdná, zabezpečující však kvalitní dosažení cíle. Z rozhodovacích alternativ zcela vymizela možnost výzkum nedělat, nejsme-li schopni nalézt adekvátní zabezpečení — šetření se provádí za každou cenu, i když kvalita výsledné statistické informace je velmi nízká a výsledky neodpovídají nákladům.

A. Pravděpodobnostní výběry

Pravděpodobnostní výběry jsou náročné na odbornou i organizační přípravu. Reprezentativitu zajišťují v pravděpodobnostním smyslu, tj. pomocí řízené náhody při výběru, pomocí modelu, který proces výběru plně popisuje, a pomocí důsledků tohoto modelu a postupů umožňujících odvodit vhodné a správné vzorce pro odhad populačních charakteristik. Výsledky jsou zobecnitelné se známou mírou spolehlivosti a neurčitosti, odvoditelnou z modelu i z dat samých. Žádný ze znaků nemá zajištěnou „úplnou reprezentativnost“ ve smyslu přesné kopie jeho výběrové distribuce k distribuci populační, avšak struktura odvozená ze souboru všech dat všech znaků je s dostatečnou pravděpodobností blízko struktury skutečné. V datech je tedy uložena dostatečně přesná informace o struktuře celkové populační variability šetřených proměnných.

V možnostech aplikace pravděpodobnostních výběrových uspořádání je velká rezerva sociologické metodologie. Systematickým používáním pravděpodobnostních výběrů je možno podstatně zvýšit metodologickou úroveň československých sociologických výzkumů, kvalitu získané statistické informace a sociologických závěrů. Tato metoda je běžná v celém světě a pro hlouběji koncipované sociologické projekty je považována za nutnou podmínku.

Při jednorázovém použití je pravděpodobnostní výběr zpravidla relativně nákladný. nehledě na to, že jeho úspěšnost předpokládá opravdu dobrou přípravu. Většinou je i časově náročnější (a to právě vzhledem k delší přípravě). To pochopitelně záleží na dostupnosti seznamu jednotek, či dokonce na možnosti přímé dosažitelnosti jednotek bez prostřednictví seznamu. Například výběr z městských populací, ze závodů a podniků apod. nečiní obvykle neřešitelné problémy.

Rezerva v aplikaci metody je však především v systematickém používání pravděpodobnostních výběrů, a to v opakování výběru z jedné dlouhodobě stabilní výběrové základny, která představuje vlastně předposlední výběrový stupeň. V takovém případě se i pro posloupcnost celostátních šetření koncipuje výběrové uspořádání jednou za deset či pět let a pro každou jednotlivou realizaci se provádí pouze výběr posledního stupně.

B. Kvótní výběr

Kvótní výběr je založen na intuitivně chápaném principu: zajistíme úplnou reprezentativitu rozložení určitých základních a kontrolovatelných znaků, jejichž populační distribuce je známa, a předpokládáme, že zajištění těchto kvót povede k zajištění reprezentativnosti i ze všech ostatních hledisek. Skutečná reprezentativita ovšem závisí na mnoha činitelích, jimiž jsou:

- kvalita tazatelské sítě, její poctivost a akceschopnost,
- pravidla pro výběr respondentů určená tazatelem,
- mechanismy fungující při výběru respondentů tazateli a důslednost naplňování kvót,
- mechanismy fungující u respondentů při odmítnutích.

U pravděpodobnostních šetření je známo procento odmítnutí, což může být základem pro odhad možného vychýlení, u kvótního výběru neznáme odklon od optimálního modelu.

U kvótního výběru vybírá tazatel tak dlouho, až naplní kvóty, a tudíž případné (a velmi pravděpodobné) vychýlení stále prohlubuje.

Reprezentativita ovšem závisí též na zkoumaném problému a na tom, jak silně korelují kvótní znaky s ostatními a jakou hloubku statistické informace chceme reprezentovat.

Kvótní výběry se osvědčily především tam, kde šlo o rychlé zjišťování jednoduchých statistických vlastností u jednotlivých otázek, popřípadě o jednoduché komparativní závěry z kontingenčních tabulek. Především se to týká výzkumů veřejného mínění. Základní podmínkou však jsou vhodně volené kvóty korelované s problematikou, jednoduchá instrumentace a kvalitní tazatelská síť i instrukce pro ni.

V sociologii, kde jsou formulovány komplexní analytické úlohy explanativního a komparativního typu, je použití kvótního výběru přes všechny jeho praktické výhody většinou nevhodné a nedostačující.

Zásadní metodologické stanovisko k tomuto postupu však dosud chybí. Názory prezentované v této stati jsou založeny na zkušenostech praxe výběrových šetření tak, jak jsou uváděny v odborné literatuře, i na vlastní výzkumné praxi autora. Konečné slovo nemůže být výsledkem metodologických spekulací, může být založeno pouze na podrobných komparativních studiích různých postupů založených na empirickém materiálu a přímo pro takové metodologické záměry koncipovaných.

Velice však závisí na volbě kvótních znaků pro každý jednotlivý konkrétní problémový okruh. Tato volba nemůže být prováděna pouze od metodologického stolu, ale na základě empirických zkušeností dlouhodobé praxe. I když ke konečnému závěru má metodologie ještě daleko, lze říci, že kvótní výběr reprezentativnost předem nezajišťuje ani při dodržení všech kvót. Kvótní výběr tedy nelze ani jednoznačně přijmout, ani jej jednoznačně odmítnout. Jako u většiny postupů i zde platí obecná pravda: metoda má své místo ve výzkumu, toto místo najít je však obtížné.

C. Záměrný výběr

spočívá v úmyslné selekci jednotek, které jsou předmětem šetření. Je reprezentativní vzhledem k záměru výzkumníka. Jednotky jsou vybírány tak, aby splňovaly jednoznačně určená kritéria. Proto je záměrný výběr vlastně ekvivalentní definici cílové populace, která je totožná s výběrem. I tento typ výběru má svůj metodologický význam. Při aplikaci postupu však je osobní odpovědnost výzkumníka velmi vysoká — odpovídá nejen na vhodnost určení výběru vzhledem k záměru, ale také vzhledem k pozdějším snahám o případná zobecnění. Záměrný výběr nemůžeme diskutovat obecně — každý jednotlivý případ je nutno z hlediska reprezentativity hodnotit zvlášť.

6. Ověřování reprezentativity

O ověřování reprezentativity pojednáme pouze stručně a v přehledu. Problém by zasloužil zvláštní stať.

Pro ověřování reprezentativnosti se používají v podstatě tyto známé postupy:

1. Statistická kontrola rozložení vybraných znaků proti cenzovým rozložením či jiným přesným statistikám. Tato metoda je použitelná jen u malého počtu znaků, jejichž výskyt není řízen výběrovým postupem. Bohužel nelze tuto metodu používat u údajů subjektivních a vztahových a na speciálních populacích, o nichž nejsou vedeny žádné objektivní statistiky.

Příklad:

- a) Dotaz na průměrný prospěch ve třídě a porovnání výpovědí studentů s třídním výkazem (v praxi bylo prokázáno v tomto případě veliké zkreslení).
 - b) Porovnání příjmových rozložení podle výpovědí respondentů s rozložením ze statistických publikací.
2. Statistická kontrola rozložení vybraných znaků proti rozložení stejných znaků u jiných výzkumů.

Příklad:

Proměnné charakterizující demografické znaky i některé postojové znaky, které byly zjišťovány na náhodně vybraném souboru čtenářů *Mladého světa*, byly kontrolovány proti stejným proměnným jiného výběrového souboru pokrývajícího celou dospělou populaci ČSR; při druhém výzkumu byly kladeny i otázky: čtete Mladý svět? Porovnání obou souborů osob, které čtou MS prokázalo velice podobné distribuce demografických proměnných, přesto, že výzkum MS byl prováděn poštovním dotazníkem, zatímco v druhém výzkumu šlo o použití tazatelů, kteří opakovali návštěvu až třikrát, aby zastihli respondenta. Oba výběry byly pravděpodobnostní, asi s jedním časovým posunem.

3. Ověřování všeobecně platných faktů, které byly již dříve plně prokázány — ověřujeme průměry, procenta, distribuce, ale hlavně vztahy mezi proměnnými.

Příklad:

V průmyslové sociologii ověřujeme například strukturu odpovědí týkajících se popisu vlastní práce (odraz práce ve vědomí dělníka) pomocí faktorové analýzy.

4. Ověřování logických a očekávaných výsledků, jež mají jasnou interpretaci (procenta, průměry i korelace).

Příklad:

Očekáváme, že čtenářky *Mladého světa* se zajímají o módu daleko více než čtenáři; také u rubriky Sally a jiných rubrik očekáváme vyšší zájem u dívek než u chlapců či starších žen; zájem mladých lidí o diskotéku, stejně jako zájem starších o dechovku též patří k takovýmto výsledkům.

5. Ověřování výsledků dvou nebo více výzkumů na průniku výzkumných cílových populací, tj. na srovnatelných částech těchto výzkumů.

Příklad:

Hodnotový systém (jeho hierarchizace u mladých dělníků) — jeden výzkum se týkal mládeže od 15 do 30 let, druhý výzkum dělníků. Průnik se týkal dělnické mládeže do 30 let.

6. Nezávislé opakování výzkumu dvěma různými metodami. Nezávislé opakování tu znamená, že jeden projekt nesmí ovlivňovat druhý, a to ani při plánování, ani při praktickém vybírání.

Příklad:

- a) Stejný výzkum je proveden kvótním a zároveň pravděpodobnostním výběrem.
- b) Jsou porovnávána dvě různá uspořádání pravděpodobnostního výběru.
- c) Výzkum je prováděn dvěma různými institucemi, které organizují různé tazatelské sítě. Tímto způsobem lze prověřovat i dva kvótní výběry.

Reprezentativitu neověřují dvě intuitivně se nabízející metody:

1. Náhodný rozklad výběrového souboru na dvě části a jejich komparace (a to ani u apriorního rozkladu tazatelské sítě na dvě části — před vlastním výběrem respondentů). Je zřejmé, že máme-li soubor dat, pak jeho náhodný rozklad musí vést ke srovnatelným a v rámci náhodnosti rozkladu velice podobným distribucím. Statistické testy rozdílnosti distribucí tu ověřují správnost náhodnosti rozdělení souboru, nikoli reprezentativitu. Existuje-li vychýlení, pak se při srovnání eliminuje (viz závěr čtvrté části statí).
2. Nezávislé opakování téhož postupu dvakrát či vícekrát za sebou vede k ověření stability metody výběru, nikoli reprezentativity. Ověřuje míru neurčitosti, nikoli

vychýlení. Je-li totiž vychýlení způsobeno metodou výběru, i zde se při srovnání výsledků eliminuje.

Ve všech metodách však hrají roli různé další faktory, jako je časový posuv, řazení otázek v pořadí a jejich kontext a podobně. Empirické ověřování reprezentativity je tedy komplexní a velmi složitá úloha. Nejlepší cestou je důkladná teoretická prověrka modelu, tj. teoretické zajištění reprezentativity metody, postupu a nevychýlenosti uspořádání, a poté důkladné předvýzkumné prověření praktických aspektů, které by mohly reprezentativitu porušovat. Nejvhodnějším způsobem zajištění kvality výběrových šetření je systematická kumulace zkušeností a statistických vyhodnocení z jednotlivých akcí.

7. Reprezentativnost vzhledem k několika cílovým populacím

V sociologickém výzkumu se často zajímáme o několik populací současně, data mají sloužit k několika cílům, k různým typům zobecnění.

Uveďme si příklady:

a) Výběr návštěvníků pražských koncertů; cílové populace:

1. návštěvníci koncertů bydlicí v Praze;
2. soubor všech návštěv koncertů za celou sezónu.

b) Výběr čtenářů časopisu; cílové populace:

1. populace čtenářů;
2. populace prodaných čísel.

c) Výběr školních dětí; cílové populace:

1. populace školních dětí;
2. populace školních tříd;
3. populace zájmových kroužků;
4. populace PO.

d) Výběr dělnických kolektivů; cílové populace:

1. populace dělníků ČSSR;
2. populace dělnických kolektivů v ČSSR;
3. populace dělníků v ČSR;
4. populace dělníků v SSR;
5. populace dělnických kolektivů v ČSR;
6. populace dělnických kolektivů v SSR.

e) Výběr dospělého obyvatelstva ČSR; cílové populace:

1. dospělé obyvatelstvo ČSR;
2. muži věku 55—60 let;
3. ženy věku 50—57 et;
4. obyvatelstvo Prahy.

f) Výběr domácností ČSSR; cílové populace:

1. dospělé obyvatelstvo ČSSR;
2. populace domácností ČSSR.

Při každém z uvedených případů musíme předem pečlivě rozvážit, jak bude dvojí či vícenásobná reprezentativita zajištěna jednak pomocí vhodného výběrového uspořádání a jednak pomocí sběru či registrace nutné informace. Převody výsledků na jednotlivé populace se provádějí vážením, agregací nebo obojím.

Z těchto úvah plyne nutnost dobré přípravy výzkumu a nutnost vypořádat se s těmito úlohami už ve výzkumném projektu.

8. Závěr

Pojem reprezentativity výběrových dat je velmi komplexní a jeho diskuse ukázala na hlavní aspekty, které nutno mít na paměti v každém sociologickém výzkumu, a to už při přípravě výzkumného projektu. Praktické určování stupně reprezentativity je náročným a ne vždy řešitelným úkolem, obzvláště pokud jde o odhad vychýlení (či systematických zkreslení) způsobených metodou, odmítnutím a podobně. V této stati nebylo možno podrobně uvést metody, které odhady zpřesňují, a tím zvyšují stupeň reprezentativnosti, ani metody, které převádějí data tak, aby odpovídala různým cílovým populacím, či metody, které eliminují vychýlení. Nebylo možné ani podrobně diskutovat vztah mezi neurčitostí a vychýlením, které jsou mnohdy vzájemně svázány; například malé vychýlení (prakticky zanedbatelné) může vést k podstatnému snížení neurčitosti apod. Zde můžeme takové metody pouze vyjmenovat:

1. vážení výsledků pro jednotlivé skupiny (pro soubory);
2. vážení výsledků pro jednotlivé respondenty;
3. vynechávání případů;
4. regresní odhadové postupy;
5. poměrové odhadové postupy;
6. korekční formule;
7. sběr speciální podpůrné informace;
8. speciální výběrové kroky v plánu šetření.

Obtíže při hodnocení reprezentativity plynou především z množství různých faktorů, které ji v praxi výzkumu porušují. Při zajišťování reprezentativity mají výhodu ty instituce, které pracují rutinně se stabilní tazatelskou sítí, popřípadě se speciálním oddělením pro výběrová šetření. Neustálé metodologické vyhodnocování všech výzkumů vede k postupnému zvyšování reprezentativity, neboť umožňuje postupně nalézt faktory, které reprezentativitu narušují, umožňuje prakticky a empiricky určit optimální výběrové velikosti pro daný problém a nalézt optimální výběrové uspořádání z hlediska praktičnosti, ekonomičnosti i přesnosti.

Reprezentativitu ve výzkumech nežádka opomíjíme, respektive ji chápeme laicky, jako reprezentativitu kvótního výběru. Tato stať chce přispět k tomu, aby v našem výzkumu pojem reprezentativity získal místo, které mu náleží, a aby byl chápán v celé své komplexnosti. Výzkumník musí nutně zvažovat nikoli vnější aspekty, ale podstatu kvality či nekvalitnosti svých dat. Předkládaná výzkumná zpráva, má-li mít nějaký praktický význam, ať už pro rozvoj teorie, či pro praktická rozhodnutí na jakékoli úrovni, musí být založena na kvalitních datech. Tuto odpovědnost, tj. odpovědnost za platnost závěrů, nikdo z autora zprávy nesejme. Výzkumník je odpovědný za vynaložené prostředky, které poskytuje společnost; zde nestačí subjektivistické a laické prohlášení, že „data jsou reprezentativní“. Proto vznikla celá samostatná větev matematické statistiky, která o výběrových šetřeních teoreticky i prakticky pojednává, aby se zjišťování statistických aspektů reprezentativity

(9) Ve velkých institucích, které pravidelně provádějí společenskovední výzkumy či jiná výběrová šetření, vznikají samostatné útvary, které mají za úkol pouze připravovat výběrová uspořádání pro jednotlivé projekty a provádějí i konkrétní výběry jednotek posledního stupně uspořádání. Takováto soustředěná práce vede po několika letech k zajištění reprezentativity

na nejvyšším možném stupni a postupně též k neustále ekonomičtějším a přesnějším uspořádáním. Pravidelný sběr metodologické informace z výzkumů a o výzkumech pak vede i k optimalizaci jednotlivých úloh, neboť obsahová náplň úlohy určuje typ variability, který – je-li znám – může být ve výběrovém uspořádání reflektován.

mohlo provádět profesionálně(9) a správně. Nerespektování a podceňování jednotlivých součástí vede ve svých důsledcích k velmi nepříjemným závěrům.

Literatura

- Cochran, W. G.: *Sampling Techniques*. New York, John Wiley & Sons 1953.
Hájek, J.: *Teorie pravděpodobnostního výběru s aplikacemi na výběrová šetření*. Praha, NČSAV 1960.
Hansen, M. H. — Hurwitz, W. N. — Madow, W. G.: *Sample Survey Methods and Theory, Vol. I, Vol. II*. New York, John Wiley & Sons 1953.
Kish, L.: *Survey Sampling*. New York, John Wiley & Sons 1965.
Řehák, J.: *Poznámky k analýze sociologických dat*. Praha, Sociologický časopis 4/1971; s. 425—434.
Šljapentoch, V. E.: *Problemy reprezentativnosti sociologičeskoj informacii*. Moskva, Statistika 1976.
Yates, F.: *Vyboročnyj metod v perepisjach i obsledovanijach*. Moskva, Statistika 1965.
Zich, F.: *Sociologický výzkum*. Praha, Svoboda 1976.

Резюме

Ржегак Я.: К понятию «репрезентативность» в социологических исследованиях

В статье ключевым понятием является «репрезентативность» выборочных исследований. Это понятие является предметом дискуссии в семи частях работы:

1. Отношение выборочной и основной совокупности данных
2. Выборочный подход и выборочная совокупность данных
3. Определение понятия «репрезентативность»
4. Два основных показателя репрезентативности: неопределенность оценки и смещение
5. Типы выборки и их репрезентативность
6. Проверка репрезентативности
7. Репрезентативность с учетом нескольких популяций

В работе проанализированы различные аспекты репрезентативности:

- с точки зрения содержания;
- с учетом социологической задачи;
- с учетом покрытия инструмента;
- с учетом типа статистических свойств;
- степень покрытия целевой популяции.

Репрезентативность распространяется на оба звена выбора, на выборочный план и на совокупность данных (реализацию выбора), которые нельзя отделить друг от друга. Она определяется как отношение выборочных данных к целевой популяции: «Репрезентативностью является степень перенесения информации, содержащейся в выборочных данных, на целевую популяцию или ее выборочную часть. Посредником в этом обобщении является статистическая модель выборочной организации и знание отклонения предполагаемой модели от ее реализации. Охватывает степень покрытия роли с точки зрения содержания (количество и роль тех знаков, информация которых обобщима), степень покрытия статистических свойств данных (свойства статистических измерений, дистрибуции и общей структуры данных) и масштаб покрытия популяции (включение всех типов единиц популяции в выборку). Это — степень регистрации популяционной изменчивости размещения всех переменных величин исследовательского инструмента с помощью данной, реализованной выборки и данной модели выборочного плана.»

Операционализация понятия «репрезентативность» проводится посредством двух основных статистических указателей точности: неопределенность оценки (стандартная ошибка оценки) и смещение. В статье обсуждаются источники обоих этих аспектов и возможности того, как их преодолевать.

Вкратце приведен список методов проверки репрезентативности и проблематика репрезентативности нескольких целевых популяций путем одного варианта выборочного плана.

Отдельные понятия и аспекты документированы и иллюстрированы посредством примеров конкретной практики социологических исследований.

Summary

J. Řehák: The Concept of Representativity in Sociological Research

The paper deals with the representativity of sampling surveys. The concept is discussed in seven parts of the paper:

1. Relation between the target population and the sample
2. Sample design and the sample
3. Definition of the concept of "representativity"
4. Two basic indicators of representativity: uncertainty of the estimate and bias
5. Types of sampling and their representativity
6. The checking representativity
7. Representativity with respect to several target populations

Various aspects of the concept are analyzed:

- conceptual representativity
- with respect to the given research task
- with respect to the coverage of research instrument
- with respect to the type of statistical properties investigated
- degree of coverage of the target population

Representativity is related to both components of sampling: sample design and the sample (collection of the data, realization of the sample design) that cannot be separated. It is defined as relation between the sample and the target population. "Representativity is a degree of generalizability (transferability) of information contained in the sample on the target population or its given part; the generalizability is conditioned by the statistical model of sample design and by the knowledge of the discrepancies between the model and its empirical realization. It involves the degree of content-coverage of the research task (number and role of variables whose information is generalizable), the degree of coverage of statistical properties of collected data (properties of statistical measures, distributions and whole structure of data) and the coverage of the target population (the inclusion of all the unit types of target population into the sample). It is the degree of coverage of the overall populational variability which concerns all the variables contained in research instrumentation by the given sample and its model (sample design).

Operationalization of the concept is made with the help of two basic statistical indicators of precision: uncertainty of the estimate and its bias.

The paper discusses sources of these two aspects and also the ways to overcome them.

The basic types of sampling: probabilistic, quota and purposive sampling are evaluated with respect to the representativity. The list of methods for checking representativity is shortly introduced as well as the problem (frequent in practice) of representing several target populations by one sample.

The concepts of the paper are documented and illustrated by examples taken from social surveys and from practical field work.