

Inhaltsverzeichnis

Vorwort und Zusammenfassung	5
1 Einleitung	7
2 Stichprobenrahmen	7
3 Stichprobenplan	9
3.1 Allgemeines	9
3.2 Methode strata.LH zur Optimierung einer Stichprobengrösse	10
3.3 Nationaler Stichprobenplan	10
3.4 Konstruktion der Take-All-Schicht	11
3.5 Regionaler Stichprobenplan	13
3.6 Erwartete Genauigkeiten auf Niveau Arbeitsstätten	14
3.7 Kombiniertes, finaler Stichprobenplan	16
4 Ziehung der Stichprobe	17
5 Anpassung des Rahmens und der Stichprobe mittels Reporting	18
6 Konstruktion des Rahmens für die Extrapolation	22
6.1 Aktualisierung des Universums mit der STATENT	23
6.2 Behandlung der Antwortausfälle	24
6.3 Einsetzungen fürs Profiling und der Take-All-Schicht	25
6.4 Anpassung der Antwortausfälle mit einer Anpassung der Gewichte	26
6.5 Anpassung der Antwortausfälle für die Beschäftigtenvariablen	26
6.6 Anpassung der Antwortausfälle für die Variablen zur Anzahl offenen Stellen	28
7 Kalibrierung und Schätzverfahren	29
7.1 Kalibrierungsvariablen	29
7.2 Robustifizierung der Gewichte	30
7.3 Kalibrierung der Gewichte für die Beschäftigtenvariablen	30
7.4 Varianzschätzung für die Beschäftigtenvariablen	31
7.5 Kalibrierung der Gewichte für die Variablen zur Anzahl offenen Stellen	32
7.6 Varianzschätzung für die Variablen zur Anzahl offenen Stellen	32
7.7 Resultate	34
8 Schlussfolgerungen	36
Anhang	39
A Beschrieb Variablen	39
A.1 Definitionen Betriebstypen	39
A.2 Definitionen NOGA	39
Literatur	40
Methodenberichte der Sektion Statistische Methoden des BFS	41

Vorwort

Die Sektion Statistische Methoden METH wurde mit der Revision der statistischen Methoden der Beschäftigungsstatistik BESTA im Jahr 2015 sowie der Erneuerung der Stichprobe im Jahr 2018 beauftragt. Der Hauptteil dieser Arbeiten wurde in einer ersten Phase vom ehemaligen Mitarbeiter von METH, Jean-Marc Nicoletti, und später von Jann Poterat durchgeführt. Tatkräftig unterstützt wurden beide während dieser Zeit von Daniel Assoulin.

In der Sektion Konjunkturerhebung KE, welche zuständig für die Durchführung der Erhebung ist, waren Francis Saucy, Rick Trap, Nathalie Nünlist, Sophie Schmassmann und Alberto Columbano stets gute Ansprechpartner bei Fragen und für Diskussionen. Die Autoren danken ihnen herzlich für die gute und effiziente Zusammenarbeit.

Ein Dank geht auch an Jean-Pierre Renfer und Daniel Kilchmann von der Sektion Statistische Methoden für die fachliche Unterstützung und die aufmerksame Lektüre des Berichts.

Zusammenfassung

In diesem Methodenbericht werden die statistischen Methoden vom Stichprobenplan bis zur Hochrechnung der Beschäftigungsstatistik BESTA beschrieben. Der Fokus wird dabei auf die Erhebung des zweiten Quartals 2015 (Revision der statistischen Methoden) und jene des ersten Quartals 2018 (Neuziehung der Stichprobe) gelegt.

Der nationale Stichprobenplan der BESTA ist so konzipiert, dass die erwünschte Genauigkeit für Wirtschaftsbranchen NOGA und für die Grossregionen eingehalten wird. Der nationale Plan führte im 2018 auf eine erwartete Brutto-Stichprobengrösse von rund 16'000 Unternehmen. Den regionalen statistischen Ämtern wird die Gelegenheit gegeben, ihre Stichprobe auf ihrem Gebiet zu vergrössern, damit für ihr Gebiet eine bestimmte Genauigkeit erreicht wird. Dies führte im 2018 insgesamt auf eine erwartete Brutto-Stichprobengrösse von rund 18'800 Unternehmen.

Bis vor der Revision 2015 wurden in der BESTA Arbeitsstätten gezogen und befragt. Mit der Revision 2015 wurde die BESTA in das Stichprobenverwaltungssystem des BFS integriert, indem Unternehmen bzw. Verwaltungseinheiten gezogen werden, die Auswertungen geschehen jedoch weiterhin auf Niveau Arbeitsstätte.

Im Bericht wird die «Reporting-Prozedur» beschrieben, welche in jedem Quartal angewendet wird, um die Änderungen in der Grundgesamtheit der Unternehmen und so auch in der Stichprobe berücksichtigen zu können.

Aufgrund ihres starken Einflusses auf die Resultate, wird den ganz grossen Unternehmen während jeder Phase der Erhebung eine spezielle Aufmerksamkeit geschenkt. Dies tangiert insbesondere auch die in diesem Bericht beschriebenen statistischen Methoden in jeder Etappe vom Stichprobenplan bis zur Schätzung.

1 Einleitung

Die Beschäftigungsstatistik BESTA liefert quartalsweise Schätzungen der Anzahl Beschäftigten der Arbeitsstätten auf Niveau Schweiz, nach Wirtschaftsbranchen NOGA¹ und für die Grossregionen. Der nationale Stichprobenplan ist so konzipiert, dass die erwünschte Genauigkeit für diese erwähnten Niveaus eingehalten wird. Den regionalen Statistischen Ämtern wird die Gelegenheit gegeben, die Stichprobe auf ihrem Gebiet zu vergrössern, damit für ihr Gebiet eine gewisse Genauigkeit erreicht wird. Bis 2015 wurden in der BESTA Arbeitsstätten gezogen und befragt, mit der Revision 2015 wurde die BESTA in das Stichprobenverwaltungssystem (SVS) des BFS integriert. In diesem System entsprechen die Ziehungseinheiten des privaten Sektors den Unternehmen, im öffentlichen Sektor sind es die Verwaltungseinheiten. So wurde aus der BESTA neu eine Cluster- bzw. Klumpenstichprobe, bei der in einem ersten Schritt Unternehmen (und Verwaltungseinheiten) gezogen werden. Zu jedem im SVS gezogenen Unternehmen (oder Verwaltungseinheit) werden auch sämtliche Arbeitsstätten in die Stichprobe aufgenommen. Die Schätzungen erfolgen weiterhin auf Niveau Arbeitsstätten. Für Unternehmen, welche Bestandteil des Profiling sind, wird die Beschäftigteninformation aus den quartalsweisen Profiling-Erhebungen in die BESTA integriert.

Der Fokus dieses Methodenberichts liegt auf den im Zuge der Revision 2015 eingeführten Methoden und allfälligen leichten Anpassungen im Jahr 2018. Auf Methoden, welche aus der Revision 2007 übernommen wurden, und in den Methodenberichten (OFS: Renaud A., 2008) und (OFS: Renaud A., Panchard Ch., Potterat J., 2008) beschrieben sind, wird nur kurz eingegangen.

2 Stichprobenrahmen

Mit der Revision 2015 sollte bei der Definition der BESTA-Grundgesamtheit erstmals die Umstellung bzw. die Erweiterung des Unternehmens- und Beschäftigtenuniversums gemäss der STATENT² berücksichtigt werden. In der STATENT, die auf den Registern der AHV-Ausgleichskassen basiert, wird eine Einheit statistisch erfasst, sobald das Unternehmen für ihre Beschäftigten AHV-Beiträge bezahlt. Alle wirtschaftlichen Akteure (natürliche oder juristische Personen), die AHV-Beiträge ab der Einkommensschwelle von jährlich 2300 CHF abrechnen, werden in der STATENT als produktive Einheiten («Unternehmen») definiert. Früher (in den Betriebszählungen und in der BESTA vor der Revision) wurden alle Unternehmen berücksichtigt, die während mindestens 20 Stunden pro Woche tätig waren, und alle Beschäftigten, die mehr als 6 Stunden pro Woche arbeiteten. Mit dieser massiven Erweiterung werden neu Kleinstunternehmen und geringfügige Beschäftigte erfasst, die bisher weder in der BESTA noch in einer anderen amtlichen Statistik abgebildet wurden.

Die Grundgesamtheit der BESTA besteht aus allen Arbeitsstätten (AST) mit Beschäftigten gemäss neuer Definition im zweiten und dritten Wirtschaftssektor, was den NOGAs 5 bis 96 entspricht. Der Stichprobenrahmen der BESTA auf dem ersten Niveau, den Klumpen, besteht im privaten Sektor aus allen Unternehmen im Betriebs- und Unternehmensregister BUR, welche mindestens eine Arbeitsstätte haben, die das obige Kriterium erfüllen. Im BUR wird der Identifikator der Unternehmen als `entid` bezeichnet.

Für die öffentliche Verwaltung gibt es im Stichprobenverwaltungssystem eine etwas andere Definition der Ziehungseinheit. Eine öffentliche Verwaltungseinheit kann als Gruppe von mehreren zusammengehörenden Einheiten der öffentlichen Verwaltung (mehrere `entid`) betrachtet

¹für «Nomenclature générale des activités économiques» (Allgemeine Systematik der Wirtschaftszweige).

²für «Statistik der Unternehmensstruktur» (Nachfolger der Betriebszählung).

werden³. Daher sind im SVS die Identifikatoren der öffentlichen Verwaltungseinheit die Gruppen (gruppebur), bestehend aus dem Bund, den Kantonen, den Gemeinden bzw. Bezirksgemeinden. Die Unterscheidung, ob eine Einheit im BUR öffentlich oder privat ist, erfolgt mit der Rechtsform oder dem Betriebstyp (betyp). Für den Stichprobenrahmen und den Plan der BESTA werden die beiden Identifikatoren der Unternehmen und der öffentlichen Verwaltungen wie folgt zu einem gemeinsamen Identifikator `clusterid` kombiniert.

$$\text{clusterid} = \begin{cases} \text{entid, falls privates oder öffentliches Unternehmen mit } \text{betyp} \in \{\text{E13, E24, E27}\} \\ \text{gruppebur, falls öffentliche Verwaltung mit } \text{betyp} \in \{\text{E20, E21, E22, E23}\} \end{cases}$$

Die Liste aller existierenden Betriebstypen ist im Anhang in der Tabelle 16 erläutert. In den nachfolgenden Abschnitten und Kapiteln werden die durch `clusterid` identifizierten Einheiten vereinfachend als «Unternehmen» bezeichnet. Wie erwähnt sind die Auswertungseinheiten die Arbeitsstätten. Eine Spezialbehandlung erhalten (Mehrbetriebs-) Unternehmen aus dem ersten Sektor, im BUR mit `betyp` $\in \{\text{E01}\}$ und `NOGA` < 5 gekennzeichnet, die sowohl Arbeitsstätten im ersten wie zusätzlich auch im zweiten oder dritten Sektor haben. Diese Unternehmen werden auch in die Grundgesamtheit aufgenommen, jedoch werden die Arbeitsstätten aus dem ersten Sektor ausgeschlossen und mit den übrigen Arbeitsstätten wird die `NOGA` und der Kanton neu berechnet und der Betriebstyp auf E13 gesetzt. Die Berechnung der neuen `NOGA` und des Kantons des (restlichen) Unternehmens erfolgt mit Hilfe folgender Regel:

Die NOGA (oder der Kanton) der Arbeitsstätten mit der gemäss BUR grössten Anzahl Beschäftigten (zuerst Vollzeitäquivalente, danach Anzahl Beschäftigte) bestimmt die NOGA (Kanton) des Unternehmens. Falls es kein eindeutiges Maximum gibt, bestimmt der Hauptsitz des Unternehmens die beiden Merkmale.

Eine ähnliche Situation gibt es auch, wenn das Mehrbetriebsunternehmen bereits im zweiten oder dritten Sektor ist, aber noch Arbeitsstätten im ersten Sektor hat, welche ausgeschlossen werden. Die `NOGA` des Unternehmens wird durch Ausschluss dieser Arbeitsstätten nicht geändert, jedoch gibt es ein neues, für die Schichtung relevantes, Unternehmens-Beschäftigungstotal (im BFS oft als `betot` bezeichnet) und zudem kann sich der Kanton des Unternehmens ändern.

Die Tabelle 1 zeigt die wichtigsten BUR-Variablen zur Konstruktion des Stichprobenrahmens und -plans.

Tabelle 1 Variablen zur Konstruktion des Stichprobenrahmens und -plans.

Variable	Definition
<code>clusterid</code>	Identifikator Stichprobeneinheit.
<code>noga082_50_cadre_cid</code>	Wirtschaftsaktivität <code>NOGA</code> der <code>clusterid</code> gruppiert gemäss Klassifikation BFS-50.
<code>betot_cid_pop</code>	Anzahl Beschäftigte Total je <code>clusterid</code> .
<code>ept_cid_pop</code>	Anzahl Vollzeitäquivalente je <code>clusterid</code> .
<code>prof_prov2</code>	Identifikator (Flagge) mit der Angabe, ob die Einheit zum Profiling gehört (=1), oder nicht (=0).
<code>pro_name</code>	Name der Profilinggruppe
<code>betyp</code>	Betriebstyp der <code>clusterid</code> .
<code>autokz2</code>	Kanton der <code>clusterid</code> . Gemeinden mit Stichprobenerhöhung separat: zs für Stadt Zürich (ab 2008), wi für Stadt Winterthur (ab 2018)

Die Tabelle 17 im Anhang zeigt die Bedeutung der `NOGA` mit ihrer Gruppierung (`NOGA-BFS50`).

³Dies war bis Ende 2018 der Fall. Mit der Einführung einer neuen BUR-Informatik-Applikation (SBER) wird ab August 2019 auch im SVS jede Gruppe nur noch aus einer einzigen `entid` bestehen.

3 Stichprobenplan

3.1 Allgemeines

Die Ausarbeitung des Stichprobenplans basiert auf folgenden Grundsätzen:

- Die Präzisionsvorgaben zur Festlegung und Bestimmung der Stichprobengrösse erfolgen für Schätzungen auf Niveau Unternehmen.
- Verifikation der effektiv zu erwartenden Präzisionen für Schätzungen auf Niveau Arbeitsstätte.

Der Stichprobenplan für die BESTA wird in zwei Schritten entworfen. Zuerst wird ein nationaler Plan erstellt, welcher die Bedürfnisse des BFS abdeckt. Für die Regionen (Kantone oder Städte), welche die Stichprobe auf ihrem Gebiet auf eigene Kosten vergrössern möchten, wird ein regionaler Plan erstellt, wobei hierfür nur die Unternehmen mit Arbeitsstätten in diesen entsprechenden Gebieten betrachtet werden.

Die Einschlusswahrscheinlichkeit des Unternehmens i für den nationalen Plan wird mit $\pi_{\text{nat},i} > 0$ bezeichnet und auf Niveau `clusterid` (erste Stufe) berechnet. Die Stichprobe wird im SVS aufgrund eines Poisson-Stichprobenplans gezogen. Die Ziehungswahrscheinlichkeiten wurden aber wie im Fall einer geschichteten einfachen Zufallsstichprobe wie folgt festgelegt:

- Schichtung nach NOGA-BFS50 gekreuzt mit Grössenklassen
- Zielvorgabe von $CV = 4\%$ für eine Genauigkeit der Schätzung der Anzahl Beschäftigten auf Niveau NOGA-BFS50
- Integration des Profilings und der grossen Unternehmen als voll erhobene Schicht.

Die Einschlusswahrscheinlichkeit der Arbeitsstätten j (auf zweiter Stufe), gegeben das Unternehmen i wurde gezogen, $\pi_{\text{nat},j|i}$ beträgt $\pi_{\text{nat},j|i} = 1$. Das bedeutet, dass alle Arbeitsstätten in einem gezogenen Unternehmen erhoben werden.

Eine analoge Prozedur wird für Regionen angewandt, welche ihre Stichprobe erhöhen. Dort wird jedoch eine Genauigkeit von $CV = 3\%$ auf Niveau Wirtschaftssektor (2 und 3) vorgegeben. Die resultierende Einschlusswahrscheinlichkeit wird mit $\pi_{\text{reg},i} > 0$ bezeichnet, wobei «reg» für Region steht.

Die beiden Einschlusswahrscheinlichkeiten auf erster Stufe werden anschliessend kombiniert, indem für jedes Unternehmen das Maximum der beiden Einschlusswahrscheinlichkeiten berechnet wird. Die kombinierte Einschlusswahrscheinlichkeit $\pi_{\text{komb},i}$ ist also

$$\pi_{\text{komb},i} = \begin{cases} \max(\pi_{\text{nat},i}, \pi_{\text{reg},i}) & \text{falls } i \text{ zu einer Region mit Erhöhung gehört} \\ \pi_{\text{nat},i} & \text{sonst} \end{cases} \quad (1)$$

Mit diesen Einschlusswahrscheinlichkeiten wird dann im Stichprobenverwaltungssystem eine Poisson-Stichprobe gezogen.

Profiling und grosse Einheiten

Die Idee ist, einen Plan zu erstellen, der in jeder NOGA eine Grössenklasse mit den grossen Einheiten identifiziert, welche voll erhoben werden und zu 100% antworten müssen. Diese Schicht

wird im Folgenden als `TakeAll_R100` bezeichnet. Diese Vereinigung aller `TakeAll_R100`-Schichten soll von einer vernünftigen Grösse sein, damit die Sektion diese Fälle im Falle von Antwortausfällen prioritär behandeln und mehrfach mahnen kann, um eine Antwortrate von 100% zu erreichen.

Diese Einheiten sind vor allem die ganz grossen Unternehmen und wenn sie zu 100% antworten, tragen sie also nichts zur Varianz der Schätzung bei. Dies erlaubt im Vergleich zu einem Plan ohne solche `TakeAll_R100`-Schichten, die Stichprobengrösse zu minimieren/verkleinern, und die formulierten Präzisionsziele unter realistischen Antwortskennarien zu erfüllen.

Ebenso werden die Profiling-Unternehmen in die Stichprobe integriert, da diese dem BFS regelmässig die Beschäftigtenangaben mitteilen und so nicht zusätzlich für die BESTA befragt werden müssen. Hier ist es aber möglich, dass einige von diesen Unternehmen in einem Quartal nicht antworten können. Daher nimmt man für die Ausarbeitung des Stichprobenplans für diese Schicht eine Antwortquote von 90% an. Ein Teil der Profiling-Unternehmen wird sich mit der `TakeAll_R100`-Schicht überlappen, da die grossen Unternehmen in beiden Fällen vorkommen.

Der zufällige Teil der Stichprobe wird auf dem Teil berechnet, der weder zum Profiling noch zur `TakeAll_R100`-Schicht gehört. Mit der im R-Package «stratification» zur Verfügung gestellten Methode `strata.LH`, die im Folgenden beschrieben wird, wird versucht die entsprechende Stichprobengrösse zu minimieren.

3.2 Methode `strata.LH` zur Optimierung einer Stichprobengrösse

Die Methode zur Bestimmung der Schichtung `strata.LH` im R-Package «stratification», wurde von L.-P. Rivest und S. Baillargeon ([Baillargeon und Rivest, 2009, 2011](#)), entwickelt. Die Abkürzung LH bezieht sich auf Lavallée & Hidiroglou, welche eine Methode zur Schichtung in zwei Schritten entwickelt haben. Zuerst werden die Schichten bestimmt, anschliessend die Stichprobengrössen, wobei man sich auf eine Hilfsvariable X abstützt, welche als Approximation der Zielvariablen Y betrachtet wird. X ist in unserem Fall die Beschäftigtenvariable `betot` aus dem Stichprobenrahmen, welches mit der Beschäftigtenvariablen aus der Erhebung Y hoch korreliert ist. Die Autoren von ([Baillargeon und Rivest, 2009](#)) haben die Methode von Lavallée & Hidiroglou wie folgt verallgemeinert:

1. Bestimmung der Schichten und der Stichprobengrösse in einem Schritt mit der Funktion `strata.LH`.
2. Verschiedene Modellierungen der Zielvariablen mit Hilfe einer Schicht-Hilfsvariable X für den Fall, dass $Y \neq X$. → BESTA: nicht verwendet.
3. Berücksichtigung der erwarteten Antwortquote in der Berechnung der Schichtung und der Stichprobengrösse. → BESTA: verwendet.
4. Verschiedene Allokationsmethoden sind möglich (power allocation, Neyman allocation, proportional allocation). → BESTA: Neyman allocation.

3.3 Nationaler Stichprobenplan

Sei U die BESTA-Grundgesamtheit der Unternehmen (`cluster`) und $Y = \sum_{i \in U} y_i$ das Total der Beschäftigten gemäss Erhebung. Es wird in allen Berechnungen davon ausgegangen, dass y_i der bekannten Anzahl Beschäftigten gemäss Stichprobenrahmen x_i entspricht. Wir definieren mit $\mathcal{H} = \{1, \dots, 47\}$ die Vereinigung aller relevanten NOGA-BFS50. $Y_h = \sum_{i \in U_h} y_i$ ist das Total

der Beschäftigten in $U_h, h \in \mathcal{H}, \bigcup U_h = U$ und $\hat{Y}_h$ ist sein Horvitz-Thompson Schätzer mit einem Variationskoeffizienten

$$CV_h = \frac{\sqrt{\text{Var}(\hat{Y}_h)}}{Y_h}. \quad (2)$$

Der nationale Stichprobenplan wird auf Niveau `clusterid` (erste Stufe) mit folgendem Vorgehen bestimmt

1. Variationskoeffizient $CV_h = 4\%$ für jedes $h \in \mathcal{H}$.
2. Profiling PROF wird behandelt als voll erhobene Schicht mit erwarteten Antwortrate von 90%.
3. Bestimmung der Takeall-Schicht (`TakeAll_R100`) mit einer Nebenbedingung seiner Grösse und einer Antwortrate von 100%.
4. Berechnung von zwei optimalen Grössenklassen und einer minimalen Stichprobengrösse mit einer optimalen Allokation für die `clusterid` ausserhalb von `TakeAll_R100` und vom Profiling konform zu unseren Genauigkeitszielen ($CV_h = 4\%$ inklusive Profiling und `TakeAll_R100`).
5. Überprüfung der erwarteten Genauigkeit bei der Schätzung des Totals auf Niveau Arbeitsstätten.

3.4 Konstruktion der Take-All-Schicht

Die Take-All-Schichten spielen eine wichtige Rolle im Stichprobenplan und bei der Schätzung von Totalen, wenn die Zielvariable eine stark asymmetrische Verteilung mit grossen Werten hat (siehe dazu z.B. [Hidioglou, 1986](#)). Falls die 100% Antwortrate in dieser Schicht erreicht wird, kann das Genauigkeitsziel mit einer dementsprechenden kleineren Stichprobengrösse erreicht werden.

Wenn die Antwortrate in der Take-All Schicht im Stichprobenplan ignoriert oder überschätzt würde, wäre das Erreichen des Genauigkeitsziels bei der Schätzung in Gefahr. Bei der Erstellung des nationalen Plans wird zuerst die Schicht `TakeAll_R100` nach folgendem Muster erstellt: `TakeAll_R100` soll genügend gross sein, um eine praktikable Stichprobengrösse zu erhalten. Andererseits erfordert die prioritäre Behandlung von `TakeAll_R100` im Rahmen der Erhebung und der Datenaufbereitung (Ziel 100% Antwortquote) auch entsprechende Ressourcen. Darum darf die Grösse von `TakeAll_R100` die entsprechenden Vorgaben der Sektion nicht überschreiten.

Prozedurbeschrieb

Die Bestimmung von `TakeAll_R100` erfolgt mit einer iterativen Prozedur basierend auf der Methode `strata.LH`:

1. Setze $J=3$ die erwünschte Anzahl Grössenklassen, den Ziel-CV von $CV_h = 4\%$. Aufgrund der bisherigen BESTA-Erhebungen wurde die Antwortquote auf 90% für die Grössenklassen $j = 1, \dots, J - 1$ und auf 100% in der Klasse J (der Take-All Schicht) festgelegt.
2. Wende die Methode `strata.LH` innerhalb der $h \in \mathcal{H}$ (BFS50) an.
3. Sei $N_{h,J}$, die erhaltene Grösse in der Grössenklasse J .

4. Berechne $N_{tot,J} = \sum_{h \in \mathcal{H}} N_{h,J}$, die totale Anzahl Unternehmen in der Stichprobe in der Grössenklasse J mit einer Antwortrate von 100% (Grösse von TakeAll_R100).
5. Falls $N_{tot,J} > 600$, die gemäss Sektion maximal zulässige Grösse für TakeAll_R100, wiederhole die Prozedur mit $J = J + 1$. Sonst ist Prozedur zu Ende: die Schicht TakeAll_R100 wird definiert mit einer Grösse von $N_{tot,J}$.

Die Tabelle 2 zeigt die erhaltenen Resultate für 2015 und 2018. Es ist ersichtlich, dass $N_{tot,J}$ mit wachsendem J kleiner wird. Um eine Stichprobengrösse ≤ 600 für TakeAll_R100 zu erhalten, muss $J = 7$ gewählt werden, welches eine voll erhobene Schicht von 494 (bzw. 522) Unternehmen ergibt. Die erwartete totale Stichprobengrösse (in Tabelle 2 als $n_{s_{nat}}$ bezeichnet) wird ebenso nach jeder Iteration kontrolliert. Es ist ersichtlich, dass die totale erwartete Stichprobengrösse mit Abnahme der TakeAll_R100 zunimmt. Die Zunahme ist der Preis für die Beschränkung der Grösse von TakeAll_R100. Der Gewinn ist, dass die Präzisionsziele mit realistischen Antwortraten erreicht werden können.

Tabelle 2 Stichprobengrössen von TakeAll_R100, $N_{tot,J}$ in Abhängigkeit von J .

	2015		2018	
J	Take-All Grösse $N_{tot,J}$	Stichpr.grösse $n_{s_{nat}}$	Take-All Grösse $N_{tot,J}$	Stichpr.grösse $n_{s_{nat}}$
3	3'924	10'153	3'770	10'105
4	1'989	11'244	2'029	11'011
5	1'158	13'044	1'158	12'829
6	755	14'877	734	14'752
7	494	16'464	522	16'043

Bemerkung: Falls eine bestimmte Profilinggruppe bestehend aus mehreren Unternehmen antwortet, liefert sie in der Regel die Angaben für alle ihre Unternehmen. Es wurde daher beschlossen, alle Profilinggruppen, die mindestens ein Unternehmen in der TakeAll_R100-Schicht haben, im Nachhinein ganz in diese Schicht zu verschieben. Dies erhöhte die Anzahl Unternehmen in der TakeAll_R100-Schicht im 2018 von 522 auf 1486 Unternehmen. Falls ein Unternehmen dieser Schicht in einem Quartal nicht antwortet, wird es eingesetzt (s. Abschnitt 6.3).

Unternehmen in nicht voll erhobenen Schichten (ausserhalb TakeAll_R100, PROF)

Für die Unternehmen in nicht voll erhobenen Schichten V_h , die also weder in TakeAll_R100_h noch im Profiling (PROF_h) sind, $V_h = U_h \setminus \{\text{TakeAll_R100}_h \cup \text{PROF}_h\}$, $h \in \mathcal{H}$, wird ein angepasster Ziel-CV berechnet. Da $U_h = \text{TakeAll_R100}_h \cup (\text{PROF}_h \setminus \text{TakeAll_R100}_h) \cup V_h$ ist

$$Y_h = Y_{\text{TakeAll_R100}_h} + Y_{\text{PROF}_h \setminus \text{TakeAll_R100}_h} + Y_{V_h}, \quad (3)$$

und die Varianz von $\hat{Y}_h$ ist

$$\text{Var}(\hat{Y}_h) = \text{Var}(\hat{Y}_{\text{TakeAll_R100}_h}) + \text{Var}(\hat{Y}_{\text{PROF}_h \setminus \text{TakeAll_R100}_h}) + \text{Var}(\hat{Y}_{V_h}). \quad (4)$$

Zu bemerken ist, dass $\text{Var}(\hat{Y}_{\text{TakeAll_R100}_h}) = 0$, da TakeAll_R100 mit einer Antwortquote von 100% voll erhoben wird. Die Varianz von $\text{PROF}_h \setminus \text{TakeAll_R100}_h$, dem Anteil des Profilings, der nicht zu TakeAll_R100_h gehört, verschwindet nicht mit der Hypothese einer Antwortquote von 90%. Mit Einsetzen der Gleichungen (4) und (3) in die Formel (2) und auflösend nach $\text{CV}_{V_h} = \frac{\text{Var}(\hat{Y}_{V_h})}{Y_{V_h}}$ ergibt dies einen angepassten CV von

$$\text{CV}_{V_h} = \frac{\sqrt{Y_h^2 \text{CV}_h^2 - \text{Var}(\hat{Y}_{\text{PROF}_h \setminus \text{TakeAll_R100}_h})}}{Y_{V_h}}, \quad (5)$$

Mit (5) und $CV_h=4\%$ wird nun für jedes h ein angepasster Ziel-CV (CV_{V_h}) berechnet. Als nächstes wird in jedem $h \in \mathcal{H}$ die strata.LH Methode mit folgendem Vorgaben angewendet:

- Neyman-Allokation
- zwei Grössenklassen
- eine Antwortquote von 90%
- CV_{V_h} als Ziel-CV

Diese Etappe liefert für jedes h zwei Grössenklassen 1 und 2 mit den entsprechenden Stichprobengrössen $\tilde{n}_{V_{h,1}}$ und $\tilde{n}_{V_{h,2}}$ und den Populationstotalen $\tilde{N}_{V_{h,1}}$ und $\tilde{N}_{V_{h,2}}$. $V_{h,1}$ und $V_{h,2}$ bezeichnen die Unternehmen in Grössenklasse 1, respektive 2.

Die Einschlusswahrscheinlichkeiten des nationalen Plans werden wie folgt berechnet:

$$\pi_{nat,i} = \begin{cases} \frac{\tilde{n}_{V_{h,1}}}{\tilde{N}_{V_{h,1}}} & \text{falls die Stichprobeneinheit } i \in V_{h,1} \\ \frac{\tilde{n}_{V_{h,2}}}{\tilde{N}_{V_{h,2}}} & \text{falls die Stichprobeneinheit } i \in V_{h,2} \\ 1 & \text{falls die Stichprobeneinheit } i \in \text{PROF} \cup \text{TakeAll_R100} \end{cases} \quad (6)$$

So liefert $n_{s_{nat}} = \sum_{i \in U} \frac{1}{\pi_{nat,i}}$ die erwartete Stichprobengrösse für den nationalen Plan. Die Tabelle 2 zeigt die erwarteten Stichprobengrössen $n_{s_{nat}}$ für verschiedene Grössen der Take-All Schicht ($N_{tot,J}$).

3.5 Regionaler Stichprobenplan

Die regionalen Stichprobenpläne wurden für die Gebiete (Kantone bzw. grosse Städte) erstellt, die ihre Stichprobengrösse auf ihrem Gebiet erhöhen wollen.

Im 2015 sind es fünf und im 2018 sechs Regionen $\mathcal{K} = \{1, \dots, 6\}$, welche den Antrag gestellt haben, um eine Schätzung mit einer kontrollierten Genauigkeit auf ihrem Gebiet zu erhalten. Die Konstruktion des regionalen Plans erfolgt mit der gleichen Prozedur wie jene auf nationaler Stufe. Sei $U^* = \cup_{k \in \mathcal{K}} U_k^*$, mit U_k^* der Menge der Unternehmen der Region k . Die betrachteten Untersuchungsbereiche, für welche eine bestimmte Genauigkeit erreicht werden soll, sind die Gebiete gekreuzt mit dem Wirtschaftssektor $\mathcal{L} = \{2, 3\}$ (zweiter und dritter Sektor). Die regionalen Pläne werden auf der ersten Stufe (der Unternehmen) definiert mit einem Ziel-CV von $CV_h = 3\%$ für alle $(k, l) \in \mathcal{K} \times \mathcal{L}$. So wird allen $i \in U^*$ eine regionale Einschlusswahrscheinlichkeit $\pi_{reg,i}$ zugeordnet.

Es bleibt anzumerken, dass für den regionalen Plan, die Unternehmen bezüglich Standort (Kanton / Grossregion) und Aktivität (NOGA / Wirtschaftssektor) in wenigen Fällen umkodiert werden mussten:

- Einschränkung der Population auf Unternehmen mit Arbeitsstätten in den interessierenden Regionen (2018: ca. 240'000 Unternehmen).
- Einbezug der Kantonszugehörigkeit und Wirtschaftsaktivität aufgrund der Beschäftigten in den Arbeitsstätten: Kanton bzw. Stadt und NOGA (Sektor) mit den meisten⁴ Beschäftigten bezogen auf die interessierenden Regionen.

⁴Prozedur analog nationaler Plan: zuerst e_{pt} , bei Gleichstand: bet_{ot} , am Schluss: Hauptsitz.

Diese Umkodierung betrifft nur Mehrbetriebsunternehmen, die für den regionalen Plan berücksichtigt werden und die in mehreren Regionen bzw. mehreren NOGAs aktiv sind. Konkret wurden die folgenden vier Eigenschaften (Variablen) für folgende Anteile der Unternehmen im Rahmen geändert: NOGA 0.1%, Sektor 0.05%, Kanton/Stadt 1.1% und Grossregion 0.7%.

3.6 Erwartete Genauigkeiten auf Niveau Arbeitsstätten

Für die Erarbeitung des Stichprobenplans wurden Genauigkeiten für den nationalen und regionalen Plan vereinfachend auf Niveau Unternehmen vorgegeben. Effektiv interessieren aber die Genauigkeiten auf Niveau Arbeitsstätten, da die Resultate nur für diese Einheiten publiziert werden sollen. Darum wurden während der Entwicklung des Stichprobenplans auch die Genauigkeiten der Schätzungen auf Niveau Arbeitsstätten im Auge behalten. Diese Berechnungen der zu erwartenden Genauigkeiten wurden mit der Statistiksoftware R vorgenommen. Sie basieren auf dem Horvitz Thompson Schätzer im Fall einer zweistufigen Stichprobe. Dabei wird die erwartete Nettostichprobe auf Stufe Unternehmen als geschichtete einfache Zufallsstichprobe und die zweite Stufe als Klumpenstichprobe behandelt. Die Tabellen 3 und 4 zeigen die Genauigkeiten des nationalen Plans der Jahre 2015 und 2018 für Arbeitsstätten, einerseits auf Niveau NOGA, andererseits auf Niveau Grossregion x Wirtschaftssektor.

Nebst den Präzisionen auf Niveau NOGA (mit Zielvorgabe CV von 4% auf Niveau Unternehmen) interessieren auch die Präzisionen für die Niveaus Grossregion und Sektor x Grossregion, wobei für die beiden letzten Niveaus bei der Planerstellung keine Ziele definiert worden sind. Aus der Erfahrung der aktuellen BESTA sollte auch für die Grossregion x Sektor ungefähr eine Genauigkeit von rund 3% erreicht werden, was zu überprüfen ist. Für die Regionen mit Erhöhung wurde ein Ziel-CV von 3% für Region x Sektor vorgegeben.

Die Tabelle 5 zeigt für die erhöhenden Regionen die erwarteten Genauigkeiten nach Wirtschaftssektoren je Region für 2015 und 2018.

Die Tabellen enthalten jeweils zwei Genauigkeiten

- **CV_HT:** Der erwartete CV auf Niveau Arbeitsstätten (Clusterstichprobe) ohne Berücksichtigung einer möglichen Kalibrierung.
- **CV_rho_0.8:** Der erwartete CV nach einer Kalibrierung, wobei eine Korrelation von $\rho = 0.8$ zwischen dem Beschäftigungstotal im Stichprobenrahmen (Hilfsvariable) und der in der Erhebung beobachteten Beschäftigtenvariablen (Zielvariable) angenommen wurde. Diese Korrelation basiert auf früheren Erhebungs-/Rahmenwerten. Die Berechnung von CV_rho_0.8 erfolgte vereinfachend mittels

$$CV_rho_0.8 = \sqrt{1 - \rho^2} CV_HT. \quad (7)$$

und beruht auf der Annahme, dass die Varianz der betrachteten Totalschätzungen ohne Kalibrierung proportional zu Varianz der Zielvariablen ist, und der Tatsache, dass gemäss (Deville und Särndal, 1992) im Falle einer Kalibrierung für die Varianzberechnung die Werte der Zielvariablen (y) durch jene der Kalibrierungsresiduen (ϵ) ersetzt werden dürfen. In der einfachen linearen Regression gilt $\text{Var}(\epsilon) = (1 - \rho^2) \text{Var}(y)$. Die Verwendung dieser Beziehung führt dann schliesslich auf (7).

Tabelle 3 Erwartete Genauigkeiten, nationaler Plan, für Arbeitsstätten auf Niveau NOGA

2018						2015					
NOGA	CV_HT	CV_rho_0.8	NOGA	CV_HT	CV_rho_0.8	NOGA	CV_HT	CV_rho_0.8	NOGA	CV_HT	CV_rho_0.8
5.9	4.65%	2.79%	58.6	4.04%	2.43%	5.9	4.57%	2.74%	58.6	4.05%	2.43%
10.2	4.31%	2.59%	61	3.67%	2.20%	10.2	4.08%	2.45%	61	3.66%	2.20%
13.5	4.24%	2.55%	62.3	3.91%	2.34%	13.5	4.17%	2.50%	62.3	3.87%	2.32%
16.8	3.98%	2.39%	64	3.98%	2.39%	16.8	3.96%	2.37%	64	4.00%	2.40%
19.2	4.09%	2.45%	65	4.22%	2.53%	19.2	4.16%	2.50%	65	3.92%	2.35%
21	3.91%	2.35%	66	3.71%	2.23%	21	3.87%	2.32%	66	3.94%	2.36%
22.3	4.06%	2.44%	68	3.93%	2.36%	22.3	4.07%	2.44%	68	3.97%	2.38%
24.5	3.98%	2.39%	69	3.99%	2.39%	24.5	3.94%	2.36%	69	4.00%	2.40%
26	4.17%	2.50%	70	3.47%	2.08%	26	4.17%	2.50%	70	3.46%	2.07%
27	3.90%	2.34%	71	3.86%	2.31%	27	3.93%	2.36%	71	3.95%	2.37%
28	3.96%	2.38%	72	3.55%	2.13%	28	3.95%	2.37%	72	3.88%	2.33%
29.3	4.24%	2.54%	73.5	3.94%	2.37%	29.3	4.48%	2.69%	73.5	4.00%	2.40%
31.3	3.88%	2.33%	78	4.02%	2.41%	31.3	3.84%	2.30%	78	4.04%	2.42%
35	4.16%	2.49%	79.2	3.72%	2.23%	35	4.14%	2.49%	79.2	3.96%	2.38%
36.9	4.22%	2.53%	84	2.85%	1.71%	36.9	4.01%	2.40%	84	3.98%	2.39%
41.2	3.85%	2.31%	85	4.52%	2.71%	41.2	3.97%	2.38%	85	3.80%	2.28%
43	3.95%	2.37%	86	3.95%	2.37%	43	3.94%	2.36%	86	3.98%	2.39%
45	3.96%	2.38%	87	3.82%	2.29%	45	3.97%	2.38%	87	3.84%	2.31%
46	3.88%	2.33%	88	3.69%	2.22%	46	3.90%	2.34%	88	4.03%	2.42%
47	4.29%	2.58%	90.3	3.52%	2.11%	47	4.34%	2.60%	90.3	3.88%	2.33%
49	3.91%	2.35%	94.6	3.90%	2.34%	49	3.92%	2.35%	94.6	3.95%	2.37%
50.1	4.06%	2.44%				50.1	3.92%	2.35%			
52	3.33%	2.00%				52	3.33%	2.00%			
53	4.14%	2.48%	Mean	3.94%	2.36%	53	4.07%	2.44%	Mean	3.97%	2.38%
55	3.94%	2.36%	Min	2.85%	1.71%	55	3.93%	2.36%	Min	3.33%	2.00%
56	3.72%	2.23%	Max	4.65%	2.79%	56	3.77%	2.26%	Max	4.57%	2.74%

Tabelle 4 Erwartete Genauigkeiten, nationaler Plan, für Arbeitsstätten auf Niveau Grossregion x Sektor

Grossregion	Sektor	2018		2015	
		CV_HT	CV_rho_0.8	CV_HT	CV_rho_0.8
1	2	4.70%	2.82%	4.58%	2.75%
	3	2.44%	1.46%	2.40%	1.44%
	Total	2.18%	1.31%	2.17%	1.30%
2	2	3.43%	2.06%	3.41%	2.04%
	3	2.53%	1.52%	2.54%	1.52%
	Total	2.09%	1.25%	2.11%	1.27%
3	2	4.17%	2.50%	3.96%	2.38%
	3	3.40%	2.04%	3.27%	1.96%
	Total	2.75%	1.65%	2.62%	1.57%
4	2	5.18%	3.11%	5.01%	3.00%
	3	2.21%	1.32%	2.35%	1.41%
	Total	2.04%	1.23%	2.15%	1.29%
5	2	3.99%	2.40%	3.95%	2.37%
	3	3.72%	2.23%	3.65%	2.19%
	Total	2.89%	1.73%	2.83%	1.70%
6	2	5.15%	3.09%	5.38%	3.23%
	3	3.81%	2.29%	3.97%	2.38%
	Total	3.15%	1.89%	3.29%	1.97%
7	2	8.45%	5.07%	8.46%	5.08%
	3	6.36%	3.82%	6.66%	4.00%
	Total	5.31%	3.19%	5.59%	3.36%

Es ist ersichtlich, dass mittels CV_rho_0.8 fast alle Zielvorgaben im Allgemeinen auch auf Niveau Arbeitsstätten eingehalten oder sogar noch verbessert werden. Einzig für das Tessin bleibt der CV über der Zielvorgabe. Die Unterschiede kommen von den Mehrbetriebsunternehmen, welche Arbeitsstätten aus unterschiedlichen NOGA bzw. Regionen haben. Die Berechnung der Genauigkeit basiert auf der Hypothese einer Antwortquote von 100% für die Unternehmen in der TakeAll_R100 und von 90% für die übrigen Unternehmen (inkl. Profiling). Da die Unternehmen die zu befragenden Einheiten auf dem ersten Niveau sind und diese die Angaben zu den Arbeitsstätten liefern, wurden nur Antwortausfälle auf diesem Niveau berücksichtigt.

Tabelle 5 Erwartete Genauigkeiten, regionaler Plan, für Arbeitsstätten auf Niveau Region x Sektor

Region	Sektor	2018		2015	
		CV_HT	CV_rho_0.8	CV_HT	CV_rho_0.8
GE	2	2.92%	1.75%	3.19%	1.91%
	3	2.75%	1.65%	2.77%	1.66%
	Total	2.12%	1.27%	2.43%	1.46%
NE	2	2.94%	1.76%	2.78%	1.67%
	3	2.79%	1.68%	3.08%	1.85%
	Total	2.37%	1.42%	2.24%	1.34%
SG	2	2.95%	1.77%	2.93%	1.76%
	3	2.67%	1.60%	2.77%	1.66%
	Total	2.04%	1.22%	2.12%	1.27%
VD	2	2.87%	1.72%	2.95%	1.77%
	3	2.75%	1.65%	2.97%	1.78%
	Total	2.42%	1.45%	2.50%	1.50%
zs	2	2.93%	1.76%	3.31%	1.98%
	3	2.88%	1.73%	3.03%	1.82%
	Total	2.45%	1.47%	2.84%	1.70%
wi	2	3.13%	1.88%	-	-
	3	2.87%	1.72%	-	-
	Total	2.71%	1.63%	-	-

Dabei steht zs für Zürich Stadt, und wi für die Stadt Winterthur

3.7 Kombiniertes, finaler Stichprobenplan

Der definitive, finale Stichprobenplan ist das Resultat der Kombination der Einschlusswahrscheinlichkeiten des regionalen $\pi_{reg,i}$ und des nationalen Plans $\pi_{nat,i}$ (aus Formel 6) wie folgt:

$$\pi_{komb,i} = \begin{cases} \max(\pi_{nat,i}, \pi_{reg,i}) & \text{falls } i \in U^* \\ \pi_{nat,i} & \text{sonst} \end{cases} \quad (8)$$

Diese Wahrscheinlichkeiten werden benutzt, um im Stichprobenverwaltungssystem des BFS eine Poisson-Stichprobe zu ziehen. Die Rahmentotale sowie die erwarteten und effektiven Stichprobengrößen nach der Ziehung sind in der Tabelle 7 dargestellt. Die Tabelle 6 zeigt den Effekt der regionalen Erhöhungen auf die erwartete Bruttostichprobengröße.

Tabelle 6 Effekt der regionalen Erhöhungen auf die erwartete Bruttostichprobengröße im 2015 und 2018

Region	2015				2018			
	Rahmen	Stichprobe			Rahmen	Stichprobe		
	Anz Unt	Kombiniert	National	Differenz	Anz Unt	Kombiniert	National	Differenz
SG	33'063	1'884	1'169	715	30'610	1'699	1'126	573
VD	50'163	2'065	1'564	501	47'988	1'960	1'485	474
NE	11'274	888	333	555	10'772	788	338	450
GE	35'821	1'561	1'084	478	34'850	1'585	1'059	526
zs	40'502	1'791	1'552	239	39'480	1'825	1'511	313
wi	-	-	-	-	6'644	568	217	351
Total	170'823	8'190	5'702	2'488	170'344	8'424	5'737	2'687

Tabelle 7 Rahmen und Plan 2018

NOGA Rahmen	Anz. UNT Rahmen	Anz. UNT voll erh. Schicht	Grenze Schicht 1 zu 2	Grenze lim TA	Erwartete Stichprobengröße UNT		Anz. AST		Anz. Profiling-Einheiten		effektive Stichprobengröße	
					nationaler Plan	kombinierter Plan	Rahmen	erw. kombinierte Stichprobengr.	UNT	AST	UNT	AST
5.9	266	19	14.5	72	83	84	354	132	12	29	82	134
10.2	4'459	40	36.5	507	287	362	5'062	696	30	282	316	649
13.5	2'997	10	9.5	175	227	240	3'063	268	2	7	243	269
16.8	9'715	18	17.5	332	542	614	9'921	694	10	27	559	634
19.2	742	18	44.0	668	143	149	806	188	10	24	151	189
21	251	14	105.0	1'082	55	62	286	84	6	7	63	84
22.3	2'097	16	31.5	414	275	290	2'311	420	8	61	307	439
24.5	7'704	22	27.5	409	433	530	7'908	611	8	26	504	591
26	2'039	37	74.5	1'254	231	327	2'224	459	13	80	336	464
27	832	17	48.0	513	140	153	889	191	11	28	150	185
28	2'074	24	48.5	539	244	288	2'199	339	7	12	337	384
29.3	446	8	29.5	448	68	72	476	89	1	8	67	83
31.3	7'342	21	16.5	396	396	453	7'636	611	12	105	434	600
35	796	50	46.5	762	169	176	1'111	417	47	186	173	418
36.9	1'495	27	14.5	126	190	199	2'049	376	21	97	195	358
41.2	8'913	29	25.5	425	447	581	9'563	1'008	13	245	574	1'001
43	37'958	83	13.5	455	600	1'263	39'253	1'858	77	445	1'263	1'842
45	16'189	32	13.5	340	472	497	17'134	983	29	373	503	1'024
46	25'601	91	27.5	823	757	813	27'990	1'855	88	607	830	1'804
47	36'170	159	26.5	1'863	452	617	51'862	10'057	158	8'343	668	9'966
49	10'833	47	23.5	716	337	360	11'968	1'257	44	860	365	1'250
50.1	363	7	47.0	554	58	61	430	105	3	17	58	101
52	1'649	35	35.5	547	164	168	2'450	685	27	347	171	702
53	485	7	89.0	2'423	25	28	3'174	2'704	6	2'644	29	2'688
55	5'236	49	29.5	398	434	449	5'855	722	42	216	448	704
56	22'001	33	21.5	648	460	528	25'376	2'581	28	1'693	570	2'605
58.6	4'810	27	10.5	271	208	213	5'054	362	20	98	220	366
61	461	10	25.0	1'437	33	35	1'204	733	10	690	32	741
62.3	18'427	45	14.5	398	586	608	19'076	857	32	134	593	843
64	6'056	73	43.5	1'726	293	314	9'019	2'661	71	2'015	286	2'565
65	606	75	59.5	1'070	108	115	1'317	756	73	673	124	797
66	9'894	87	13.5	342	472	483	11'082	1'256	78	672	557	1'321
68	17'259	28	9.5	272	457	478	17'755	707	18	140	494	715
69	22'887	14	9.5	350	501	536	23'507	707	11	90	555	732
70	22'817	85	11.5	515	471	516	23'633	969	80	412	556	1'006
71	25'389	33	12.5	403	636	701	26'651	1'075	24	188	723	1'105
72	1'760	19	15.5	256	158	158	1'855	232	10	52	164	239
73.5	25'551	17	7.5	340	530	566	25'858	704	12	89	531	669
78	2'219	233	59.0	1'497	385	401	3'500	1'393	231	989	397	1'392
79.2	21'647	59	27.5	936	560	604	25'480	2'564	50	1'512	591	2'545
84	1'930	53	67.5	2'918	120	141	8'679	5'101	53	4'671	148	5'029
85	20'648	66	26.5	2'172	228	388	32'781	9'298	62	8'441	388	9'105
86	55'730	83	37.5	1'980	497	708	57'445	1'604	70	684	725	1'618
87	2'194	44	119.5	1'049	256	344	3'905	1'116	38	421	333	1'071
88	7'909	32	29.5	565	477	507	11'344	2'398	26	1'166	491	2'304
90.3	23'284	32	11.5	364	647	670	24'870	1'543	21	667	684	1'490
94.6	48'748	105	10.5	413	731	878	51'472	1'989	98	830	870	1'916
Total	548'879	2'133			16'043	18'730	626'837	67'418	1'801	41'403	18'858	66'737

Es hat schlussendlich 2133 Unternehmen in den voll erhobenen Schichten, davon sind 1556 in der Schicht der TakeAll_R100, die übrigen sind 577 kleinere Profiling-Unternehmen.

Zu den 522 aus der Tabelle 2 vom nationalen Plan wurden alle (genau 964) Unternehmen von den Profiling-Gruppen in die Schicht TakeAll_R100 hinzugefügt, welche bereits mindestens 1 Unternehmen in dieser Schicht hatten. Dies aus dem Grund, da innerhalb einer Gruppe meistens alle (bzw. keine der) Unternehmen antworten. Nur einen Teil einer Profiling-Gruppe als TakeAll_R100 zu definieren, ist daher nicht angebracht. TakeAll_R100 umfasst somit 1486 Unternehmen für den nationalen Plan. Weitere voll erhobene Unternehmen kamen vom regionalen Plan dazu.

4 Ziehung der Stichprobe

Die Bruttostichprobe S wird mit den Einschlusswahrscheinlichkeiten des kombinierten Plans gezogen. Für die Revision 2018 gibt es 118 verschiedene Einschlusswahrscheinlichkeiten. Unternehmen mit identischen Einschlusswahrscheinlichkeiten kann man für die Ziehung im SVS als eine Schicht betrachten. Im 2015, als Winterthur noch nicht erhöhte, waren es noch 115 solche Schichten. Die Anzahl Schichten ergibt sich im 2018 aus folgender Summation.

1. Alle voll erhobenen Schichten, TakeAll_R100 und jene des Profilings, werden als eine einzige Schicht betrachtet.
2. Der nationale Plan hat 47 NOGAs multipliziert mit 2 Schichten (Grössenklassen) ergibt 94.
3. Der regionale Plan hat 6 Regionen (4 Kantone, 2 Städte), die erhöhen, mit 2 Sektoren und

ebenfalls mit 2 Grössenklassen = 24. Davon wird eine Schicht zusätzlich voll erhoben, welche zur Nummer 1 gelegt wird. Es bleiben 23 nicht voll erhobene Schichten.

Im Jahr 2018 wurde beschlossen, für die BESTA ein Rotationspanel einzuführen, um die Belastung der Unternehmen gleichmässiger zu verteilen. Dabei wird eine Rotationsrate von 20% angestrebt, was bedeutet, dass unter idealen Bedingungen 20% der Stichprobe erneuert wird. Dies ist im Stichprobenverwaltungssystem möglich, indem man die zu ziehende Stichprobe in fünf Teile aufteilt, der erste Teil wird negativ mit der letzten Erhebung koordiniert, also ersetzt, vier Teile werden positiv koordiniert, also beibehalten.

Diese Anteile sind als Annäherung zu verstehen, da bei einer Erneuerung alle inaktiven Unternehmen aus Rahmen und Stichprobe entfernt werden, dafür aber alle neuen Unternehmen neu in den Rahmen aufgenommen werden. Zudem werden der Rotation durch die Grösse der Ziehungswahrscheinlichkeiten Grenzen gesetzt. Beispielsweise werden alle Unternehmen, welche im Stichprobenplan eine Ziehungswahrscheinlichkeit von 1 zugewiesen wurde, ungeachtet der Rotationsrate immer voll erhoben. Dies betrifft vor allem grössere Unternehmen.

Die Tabelle 8 zeigt für die BESTA die Grösse des Rahmens, der gezogenen Stichprobe und des Profilings global und nach Wirtschaftssektor, sowohl auf Niveau Unternehmen wie auch auf Niveau Arbeitsstätten.

Tabelle 8 Grösse des Rahmens, der Stichprobe und des Profilings in den Jahren der Erneuerungen der BESTA-Stichprobe 2018, 2015 sowie zum Vergleich 2012.

	2018			2015			2012		
	Rahmen	Sample	Profiling	Rahmen	Sample	Profiling	Rahmen	Sample	Profiling
Sektor	Niveau Ziehungseinheit (UNT)			Niveau Ziehungseinheit (UNT)			Niveau Unternehmen (UNT)		
2	90'126	5'754	288	95'887	5'776	283	78'512	9'305	237
3	458'753	13'104	1'513	479'673	13'004	1'278	304'823	16'186	1'124
Total	548'879	18'858	1'801	575'560	18'780	1'561	383'335	25'491	1'361
Sektor	Niveau Arbeitsstätte (AST)			Niveau Arbeitsstätte (AST)			Niveau Ziehungseinheit (AST)		
2	95'111	8'324	1'669	102'648	8'276	1'620	82'881	11'599	1'372
3	531'726	58'413	39'734	566'156	57'750	38'818	364'389	51'884	33'332
Total	626'837	66'737	41'403	668'804	66'026	40'438	447'270	63'483	34'704

Es ist ersichtlich, dass sich mit der Einführung der STATENT der Rahmen im 2015 wesentlich vergrössert hat. Hingegen hat sich die globale Stichprobengrösse in Arbeitsstätten seit 2012 nur unwesentlich verändert. Es werden statt anfänglich 63'500 neu knapp 67'000 Arbeitsstätten befragt, zudem gab es eine leichte Verschiebung der Stichprobe vom zweiten in den dritten Sektor. Die BESTA-Stichprobe wurde bis vor der Revision 2015 auf Niveau Arbeitsstätten gezogen, seither ist die Ziehungseinheit das Unternehmen. Die Revision 2015 bewirkte, dass die Anzahl zu befragender Einheiten auf dem ersten Niveau von 25'500 auf 18'800, oder um gut einen Viertel abgenommen hat. Dies ist im Hinblick auf die Belastung von Unternehmen als sehr positiv zu werten. Die leichte Zunahme der Stichprobe in Bezug auf die Anzahl Arbeitsstätten kann durch die Erweiterung des Profilings von 1'361 auf 1'800 Gruppen mit 34'700, respektive 41'400, Arbeitsstätten erklärt werden.

5 Anpassung des Rahmens und der Stichprobe mittels Reporting

Die BESTA ist eine Quartalerhebung, d.h. jeden dritten Monat müssen die Fragebogen neu an die Unternehmen in der Stichprobe verschickt werden. Zwischen den einzelnen Erhebungen werden im Betriebs- und Unternehmensregister BUR laufend Aktualisierungen vorgenommen. Diese betreffen einerseits den Bestand (Mutationen wie Schliessungen, Neugründungen,

Fusionen, Splits), andererseits auch die Variablen (Änderung der NOGA, Adressen, Kanton, Beschäftigte etc.). Das Ziel der BESTA ist die Produktion von Konjunktur-Schätzungen. Dabei ist es wichtig, zuverlässige quartalsweise Evolutions-Schätzungen produzieren zu können, welche nicht durch administrative Prozesse beeinflusst werden. Unter diesem Aspekt wurden demographische und Klassifikations-Änderungen nach den im Folgenden beschriebenen Prinzipien behandelt. Im Gegensatz zu konjunkturellen Veränderungen, fassen wir demographische Veränderungen und Umklassifizierungen (z.B. des NOGA-Codes) als strukturelle Veränderungen zusammen.

Für NICHT-Profiling Einheiten (Unternehmen und Arbeitsstätten):

- Strukturelle Änderungen werden im BUR nicht unbedingt zeitnah, sondern mit einer gewissen Verzögerung erfasst. Würden diese Änderungen basierend auf dem BUR-Aktualisierungsdatum in die Schätzungen übernommen, könnten diese unerwünschter Weise durch administrative Prozesse beeinflusst werden.
- Strukturelle Änderungen werden für die quartalsweisen Schätzungen weitgehend nicht berücksichtigt. Die Schätzungen basieren auf dem, gemäss Abschnitt 6, leicht angepassten Stichprobenrahmen.

Für das erste Quartal nach der Ziehung einer neuen Stichprobe wird auch die alte Stichprobe erhoben (Doppelstichprobe). Aufgrund der Rahmenbedingungen (ähnlicher Stichprobenplan, rotierendes Panel mit Rotationsrate von 20%) darf von einer starken Überlappung der beiden Stichproben ausgegangen werden. Der zusätzliche Erhebungsaufwand und die zusätzliche Belastung halten sich im Rahmen. Differenzen zwischen den Schätzungen ($\hat{y}_{new}$ vs. $\hat{y}_{old}$) können als strukturelle Änderungen interpretiert werden.

Die quartalsweisen Schätzungen basierend auf der alten Stichprobe werden dann revidiert, um diese Änderungen zu reflektieren. Die hierzu verwendete Methode ist in (OFS: Graf M., 2001) beschrieben.

Für Profiling Einheiten:

- Strukturelle Änderungen werden im BUR zeitnah erfasst.
- Strukturelle Änderungen bei Profiling Einheiten (Unternehmen und Arbeitsstätten) werden für die quartalsweisen Schätzungen berücksichtigt.

Die strukturellen Änderungen werden im oben beschriebenen Sinne durch ein als «Reporting» bezeichnetes Verfahren behandelt, welches folgend dokumentiert wird.

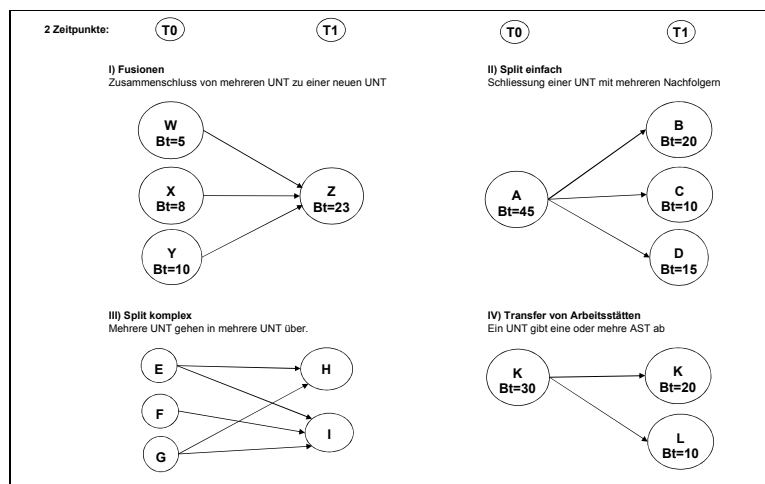
Reporting

Das Reporting hat insbesondere folgende Ziele:

- Quartalsweises Update des Stichprobenrahmens und der Stichprobe.
- Eine Berücksichtigung der Mutationen für den Versand der Fragebogen und für die Anpassung des Rahmens bei der Hochrechnung.
- Anpassung der Rahmen-Variablen, insbesondere der Schichtinformationen (betot und NOGA), bei Mutationen und dem Profiling.

Die Abbildung 1 zeigt vier Beispiele von möglichen Mutationen, die mit dem Reporting behandelt werden sollten. Dabei ist «Bt» die Anzahl Beschäftigte.

Abbildung 1 Beispiele Transfer von Unternehmen



Nicht alle Fälle kommen gleich häufig vor. Die Tabelle 9 zeigt die innerhalb von 7 Monaten registrierten Änderungen zwischen zwei Zeitpunkten im Jahr 2015 bezogen auf die Anzahl Unternehmen.

Tabelle 9 Veränderungen im BUR zwischen zwei Zeitpunkten (Intervall von 7 Monaten).

Beschrieb	Anzahl Unternehmen		
	absolut	relativ	davon Profiling
Keine Änderung	550'329	91.2%	1'486
Schliessung ohne Nachfolger	22'656	3.8%	62
Neue Einheit ohne Vorgänger	25'315	4.2%	31
Schliessung mit mind. 1 Nachfolger	2'591	0.4%	13
Neue Einheit mit genau 1 Vorgänger	2'535	0.4%	7
Neue Einheit mit mehreren Vorgängern	24	0.0%	1

Damit eine Fusion gemäss obiger Abbildung 1 zustande kommt, müssen also zwei oder mehr Unternehmen von einer Schliessung mit Nachfolger betroffen sein und der Nachfolger muss für alle derselbe sein. Komplexe Splits wurden erst bei der Analyse über eine längere Zeitdauer entdeckt und fehlen in der Tabelle 9. Der Fall von Transfers von Arbeitsstätten kommt zwar vor, wird aber aktuell im Reporting der BESTA nicht automatisch behandelt, da die Erhebung und auch die Mutationen nur auf Niveau Unternehmen betrachtet werden. Dies gab bei der Entwicklung der Methoden Anlass zu Diskussionen, insbesondere falls grosse Unternehmen betroffen waren. Um solchen Situationen und der möglichen Komplexität von Mutationen Rechnung zu tragen, wurden die Mutationen in Einzelfällen auch manuell behandelt.

Wie in der Tabelle 9 ersichtlich, ist der grösste Teil der Unternehmen (91%) nicht von Änderungen betroffen. Für die restlichen Fälle wurden aber Regeln erstellt, wie diese gehandhabt werden sollen. Dabei wird eine Unterscheidung vorgenommen, ob es sich um Fälle des Profiling handelt oder nicht. Dies aus dem Grund, da man davon ausgehen kann, dass Änderungen im Profiling meistens zeitnah im BUR vorgenommen werden, da ein regelmässiger Kontakt zwischen dem BUR und den Unternehmen stattfindet und Änderungen für Anzahl Beschäftigte, NOGA oder Anzahl Arbeitsstätten direkt mitgeteilt werden. Diese Änderungen sollen für das Profiling so direkt

in das Update des Stichprobenrahmens und der Stichprobe einfließen.

Der an die Mutationen angepasste Stichprobenrahmen wird als Hochrechnungs- oder Extrapolationsrahmen bezeichnet.

Bei den anderen Unternehmen werden Aktualisierungen meistens über Drittquellen wahrgenommen, z.B. publiziert die STATENT jährlich Ende August die Anzahl Beschäftigte, die via AHV-Daten übermittelt wurden. Diese beziehen sich im Jahr T auf den Stand am 31.12. des Jahres T-2. Ein anderes Beispiel sind Löschungen oder Gründungen von Unternehmen, die über das Schweizerische Handelsamtsblatt (SHAB) mitgeteilt werden und nur mit Verzögerung eingebunden werden.

Die Mutationen der übrigen Unternehmen ausserhalb des Profiling werden wie folgt gehandhabt:

- «Schliessung ohne Nachfolger» und «Neue Einheit ohne Vorgänger» werden im Hochrechnungsrahmen und somit für die Schätzungen nicht berücksichtigt.
- «Schliessung mit Nachfolger(n)» werden ersetzt mit den «neuen Einheiten mit Vorgängern» (Aktualisierung Stichprobe: Versand an «Nachfolger/Neue Einheit»).

Die Tabelle 10 zeigt für mehrere Transfertypen die Behandlung mehrerer Variablen für das Update im Hochrechnungsrahmen. Die Variablen sind folgende:

BetotRahmen Beschäftigungstotal $betot$ aus dem Stichprobenrahmen (gültig auch für ept)

Prof-Flag Flag, ob die Einheit im Profiling ist

Idech_neu Flag, ob die Einheit in der Stichprobe ist

poidstir Initialgewicht

inTakeAllInt_neu Flag, ob die Einheit in den voll erhobenen Schichten (TakeAll_R100) ist

straTir Schichtvariable

takoTir Flag, ob die Einheit in der voll erhobenen Schicht (inkl. Profiling) ist

Tabelle 10 Behandlung der Variablen bei Mutationen

Typ	Beschrieb	Behandlung Variablen						
		BetotRahmen	Prof-Flag	Idech_neu	poidstir	inTakeAllInt_neu	straTir	takoTir
1	Fusionen 1 zu 1 (1 Spender / 1 Nachfolger)	T0	T0	T0	T0	T0	T0	T0
2	Fusion: mind. 2 Spender für eine clusterid	sum T0	max T0	max T0	1 für profT1=1, sonst «Teilung der Gewichte»	max T0	miss	max T0
3	Split einfach (1 Spender, mehrere clusterids als Empfänger)	T1	T0	T0	T0	T0	T0	T0
4	Gemeinde-Fusionen (mehrere Spender für eine neue Gemeinde)	sum T0	max T0	max T0	1 für profT1=1, sonst «Teilung der Gewichte»	max T0	miss	max T0
5	Split komplex (mehrere Spender für mehrere neue clusterids)	T1	max T0	max T0	1 für profT1=1, sonst «Teilung der Gewichte»	max T0	miss	max T0

T0 (oder T1) heisst: Wert aus T0 (oder T1) wird übernommen / max: Maximum aller Werte / sum: Summe aller Werte

T0 heisst, dass der Wert aus dem Initialrahmen übernommen wird. Die Variable «sum T0» im Bezug auf BetotRahmen entspricht der Summe der entsprechenden Beschäftigtenzahlen sämtlicher Spenderunternehmen. Entsprechend ist «max T0» in Bezug auf idech_neu genau dann 1, wenn mindestens ein Spenderunternehmen in der Stichprobe ist. T1 bedeutet, dass man die Variable aufgrund der neuen Information zum Zeitpunkt T1 aktualisiert.

Die Berechnung der neuen Gewichte (poidstir) im Fall von Splits oder Fusionen beruht auf dem Artikel über die «Teilung der Gewichte» (von [J.-C.Deville und Lavallée, 2006](#)) und wird im Folgenden kurz beschrieben.

U_{T0}, N_{T0} bezeichnen die Unternehmenspopulation, respektive deren Anzahl Einheiten, zum Zeitpunkt $T0$ und U_{T1}, N_{T1} die entsprechenden Grössen zum Zeitpunkt $T1$. Die Unternehmen in U_{T0} werden mit j , jene in U_{T1} mit i indiziert. Analog zu ([J.-C.Deville und Lavallée, 2006](#)) definieren wir

$$\tilde{\theta}_{j,i} = \frac{\text{betot}_j \cdot 1(i \text{ Nachfolger von } j)}{\sum_{j \in U_{T0}} \text{betot}_j \cdot 1(i \text{ Nachfolger von } j)}. \quad (9)$$

Bemerkung: Definitionsgemäss gilt $\sum_{j \in U_{T0}} \tilde{\theta}_{j,i} = 1$. Basierend auf ([J.-C.Deville und Lavallée, 2006](#), Formel 2.5) lässt sich das Total der Variablen y in der Population U_{T1} mittels

$$\begin{aligned} \hat{Y}_{T1} &= \sum_{j=1}^{N_{T0}} 1(j \in S) \frac{1}{\pi_j} \sum_{i=1}^{N_{T1}} \tilde{\theta}_{j,i} y_i \\ &= \sum_{i=1}^{N_{T1}} \underbrace{\sum_{j=1}^{N_{T0}} 1(j \in S) \frac{1}{\pi_j} \tilde{\theta}_{j,i}}_{w_i^*} y_i, \end{aligned} \quad (10)$$

schätzen, wobei w_i^* das Gewicht der Einheit $i \in U_{T1}$ ist. Aus (10) folgt, dass eine Einheit $i \in U_{T1}$ genau dann in die Stichprobe aufgenommen werden muss, wenn mindestens ein Vorgänger in der Stichprobe war. Das Gewicht für $i \in U_{T1}$ ergibt sich mittels $\tilde{\theta}_{j,i}$ als gewichtete Summe der Gewichte π_j^{-1} jener Vorgänger, die sich in der Stichprobe befinden.

Beispiel Fusion: Wir nehmen an, die Unternehmen

- u_1 ; $\text{betot}_1 = 30, 1(u_1 \in S) = 1, \pi_1 = 1$
- u_2 ; $\text{betot}_2 = 15, 1(u_2 \in S) = 1, \pi_2 = \frac{1}{2}$
- u_3 ; $\text{betot}_3 = 5, 1(u_3 \in S) = 0, \pi_3 = \frac{1}{4}$

fusionieren zu u_4 . Da u_1 und u_2 in der Stichprobe waren, wird das Nachfolgeunternehmen u_4 in die Stichprobe aufgenommen ($\max\{1(u_1 \in S), 1(u_2 \in S), 1(u_3 \in S)\} = 1$). Mit (10) erhalten wir das Gewicht w_4^* von u_4 durch

$$\sum_{j=1}^3 1(u_j \in S) \frac{1}{\pi_j} \frac{\text{betot}_j}{30 + 15 + 5} = 1 \cdot \frac{30}{50} + 2 \cdot \frac{15}{50} = \frac{6}{5}.$$

Beispiel Split: Wir nehmen an, dass das Unternehmen u_1 ($\text{betot} = 30, 1(u_1 \in S) = 1, \pi_1 = 0.5$) sich in 2 Nachfolgeunternehmen u_2 und u_3 aufteilt. Da u_1 in der Stichprobe war, werden auch u_2 und u_3 in die Stichprobe aufgenommen. Die Gewichte von u_2 und u_3 entsprechen gemäss (10) dem Gewicht von u_1 ($\frac{1}{\pi_1} = 2$).

6 Konstruktion des Rahmens für die Extrapolation

Dieser Abschnitt beschreibt die Konstruktion des Rahmens für die Extrapolation auf Niveau Arbeitsstätten. Die Erhebung wird vierteljährlich durchgeführt und basiert auf dem Ziehungsrahmen auf Niveau Unternehmen (`clusterid`) zum Zeitpunkt der Ziehung $T0$ (`cadre_cid_t0`). Wie

im Kapitel 5 beschrieben, wird dieser Rahmen aufgrund der Entwicklung der Unternehmen zwischen dem Initialrahmen und dem Zeitpunkt der Erhebung T_1 via Reporting angepasst. So entsteht der angepasste Ziehungsrahmen oder Hochrechnungsrahmen auf Niveau Unternehmen (`clusterid`) in T_1 (`cadre_cid_t1`).

Im nächsten Schritt muss mit den Erhebungsdaten der Hochrechnungsrahmen auf Niveau Arbeitsstätten erstellt werden. Dies geschieht in folgenden Schritten:

1. Integration der Daten/Antworten des Profilings.
 - Die Daten des Profilings sind auf Niveau Arbeitsstätten erhoben.
 - Vergleich mit den Informationen im Rahmen `cadre_cid_t1`:
 - Die Information der Arbeitsstätte wird nur als Profiling berücksichtigt, falls ihre `clusterid` im `cadre_cid_t1` ebenfalls als Profiling-Unternehmen markiert ist.
 - Falls umgekehrt aufgrund einer Mutation die `clusterid`, für die man eine Information zu einer Profiling-Arbeitsstätte erhalten hat, nicht als Profiling im Rahmen `cadre_cid_t1`, markiert ist, wird die Arbeitsstätte nicht dem Profiling zugeordnet.
2. Einbindung der Daten ausserhalb des Profilings. Dabei werden mit EUNT, die Unternehmen mit nur einer Arbeitsstätte und mit MUNT, die Unternehmen mit mehreren Arbeitsstätten bezeichnet.
 - Die Daten ausserhalb des Profilings werden auf Niveau `clusterid` erhoben.
 - Die Unternehmensgruppen, bestehend aus mehreren Unternehmen, liefern die Daten für die ganze Gruppe aggregiert. Um die Information auf Niveau `clusterid` zu erhalten, muss eine Aufteilung mit einem Verteilschlüssel vorgenommen werden.
 - Die Daten auf Niveau `clusterid` (MUNT und EUNT) werden anschliessend ebenfalls mit den Rahmen `cadre_cid_t1` verknüpft. Falls eine `clusterid` als nicht Profiling erhoben wurde, aufgrund einer Mutation und gemäss Rahmen `cadre_cid_t1` neu zum Profiling gehört, werden diese Daten aufgrund eines Verteilschlüssels, basierend auf der Anzahl Beschäftigten, auf die Arbeitsstätten verteilt und neu ins Profiling integriert.
3. Die Behandlung der Antwortausfälle auf Niveau `clusterid` wird im Abschnitt 6.2 beschrieben.
4. Alle Antworten auf Niveau `clusterid` im Rahmen `cadre_cid_t1_nonprof`, das heisst ausserhalb des Profilings, werden mit einem Verteilschlüssel, basierend auf der Anzahl Beschäftigten, auf ihre Arbeitsstätten⁵ verteilt und anschliessend mit dem Arbeitsstätten-Rahmen des Profilings zusammengeführt. Alle Daten auf Niveau Arbeitsstätten zusammen, werden im Extrapolationsrahmen `cadre_ast_t1` festgehalten.

6.1 Aktualisierung des Universums mit der STATENT

In den ersten Jahren von 2015 bis und mit 2017 mussten die Standardmethoden leicht modifiziert werden, da beim Erstellen des Stichprobenrahmens und der Stichprobe der BESTA 2015 mit der Integration der STATENT 2012 ins BUR das Universum und auch die Beschäftigtenzahlen wesentlich verändert wurden. Viele Unternehmen und Arbeitsstätten wurden ins BUR geladen, andere wurden später deaktiviert, da sie gemäss STATENT nicht mehr aktiv sind. Um dieser

⁵Bei Unternehmen ohne Mutationen handelt es sich um die Arbeitsstätten gemäss T_0 . Im Falle von Mutationen (Splits, Fusionen) werden die Arbeitsstätten fallspezifisch definiert.

nachträglichen Anpassung Rechnung zu tragen, wurden zwei Teilmengen von Unternehmen definiert, STATENT-Konforme und Nicht-STATENT-Konforme. Letztere waren jene, die im Original-BESTA-Rahmen waren, aber mit dem Laden der STATENT deaktiviert wurden. Für diese Unternehmen wurden separate Untersuchungsbereiche für die NOGA (=100), die Region (=8) und den Wirtschaftssektor (=4) definiert. Diese Unterteilung wurde im Rahmen der Gewichtung verwendet, jedoch nicht als Untersuchungsbereich für die BESTA-Resultate.

Die Anpassung wurde nur für die Gewichtungen und Schätzungen 2015 bis 2017 verwendet. Beim Update des Rahmens und der Stichprobe im Jahr 2018 entfiel diese Spezialbehandlung wieder, da ab diesem Zeitpunkt alle Unternehmen konform mit der neusten STATENT sind.

6.2 Behandlung der Antwortausfälle

Die Behandlung der totalen Antwortausfälle geschieht auf zwei unterschiedliche Arten, je nach Schichten.

1. Fälle des Profilings und der Takeall-Schicht TakeAll_R100

- (a) Die nicht-antwortenden Arbeitsstätten des Profilings werden eingesetzt: Alle Arbeitsstätten werden mit dem Arbeitsstätten Rahmen `cadre_ast_t1` verknüpft. Die nicht-antwortenden Arbeitsstätten, welche als Profiling bezeichnet sind, werden anhand der Beschäftigten-Entwicklung zwischen zwei Quartalen eingesetzt. Dabei wird die Entwicklung zwischen dem Vorquartal (bzw. Initialrahmen) und dem aktuellen Quartal berechnet (siehe Abschnitt 6.3).
- (b) Entsprechend der Ausarbeitung des Stichprobenplans ist es für die Erreichung der Präzisionsziele zentral, in der Schicht der TakeAll_R100 eine Antwortrate von 100% zu erreichen. Dies bedeutet, dass diesen Unternehmen in der Datenerhebung und im Datenaufarbeitungsprozess höchste Priorität eingeräumt werden muss. Die Situation kann als Spezialfall eines «static adaptive surveydesign» (gemäss Schouten B., 2017) oder auch als Spezialfall von «selective editing» (gemäss Luzi O., 2007) gesehen werden.
Für Unternehmen, für die trotz allen Anstrengungen keine Antwort vorliegt, werden die Daten eingesetzt. Die entsprechende Prozedur wird im Abschnitt 6.3 beschrieben.

2. Fälle ausserhalb des Profilings und ausserhalb der Takeall-Schicht TakeAll_R100

- (a) Für diese Fälle wird eine Modellierung der Antwortausfälle auf Niveau `clusterid` vorgenommen, um angepasste Gewichte zu berechnen, mit denen dann die Beschäftigtenzahlen hochgerechnet werden können. Diese Prozedur ist in Abschnitt 6.5 beschrieben.

Für die Hochrechnung der offenen Stellen wird eine separate Gewichtung für alle `clusterid` berechnet. Hierzu wird wiederum eine Modellierung der Antwortausfälle gemacht, siehe Abschnitt 6.6.

Vor der Revision 2015 waren die Stichprobeneinheiten und die Analyseeinheiten die Arbeitsstätten. Die Behandlung der totalen Antwortausfälle wurde mit der Kalibrierung auf Niveau Arbeitsstätten in einem Schritt erledigt.

Seit 2015 ist die Ziehungseinheit die `clusterid` und nicht mehr die Arbeitsstätte. Die Analysen (im Abschnitt 7.3) weisen darauf hin, dass die Antwortausfallskorrektur auf Niveau `clusterid` die Qualität der resultierenden Schätzungen hinsichtlich Bias verbessert. Entsprechend wurde für die revidierte BESTA dieses Vorgehen gewählt.

6.3 Einsetzungen fürs Profiling und der Take-All-Schicht

Die Prozedur zur Einsetzung der nicht antwortenden Arbeitsstätten des Profiling und der nicht-Antwortenden Unternehmen der Take-All-Schicht TakeAll_R100 ausserhalb des Profiling wird mit ein paar Anpassungen aus der früheren Methode der Behandlung der grossen Einheiten «very-bigs» übernommen, wie sie im Methodenbericht von METH beschrieben wurde (siehe OFS: Renaud A., Panchard Ch., Potterat J., 2008). Nachfolgend werden die wichtigsten Etappen zur Behandlung des Profiling beschrieben.

Es werden die beiden Quartale trim=aktuelles Quartal und trim1=Vorquartal definiert. Falls es kein Vorquartal gibt (am Anfang der Erhebung oder falls keine Antwort des Vorquartals vorliegt) nimmt man trim1=Initialrahmen ($T0$). Es seien S_{trim} und S_{trim1} die Stichproben der beiden erwähnten Quartale auf Niveau AST. Mit $R_{job,trim} \subset S_{trim}$ und $R_{job,trim1} \subset S_{trim1}$ werden die antwortenden Einheiten der beiden Quartale bezeichnet. Die Entwicklung zwischen den beiden Quartalen wird unter der Benutzung von $R_{job,trim} \cap R_{job,trim1}$ pro aggregierte NOGA-BFS50 (h) wie folgt geschätzt.

$$evol_h = \frac{\sum_j EPT_{trim,j}}{\sum_j EPT_{trim1,j}}, h \in \mathcal{H}, j \in R_{job,trim} \cap R_{job,trim1} \quad (11)$$

Mit $y_{trim,j}$ wird eine der Beschäftigtenvariablen der Arbeitsstätte j im Quartal trim aus der Erhebung bezeichnet, die in folgender Liste aufgeführt wird:

- Beschäftigungstotal (tot), Vollzeitäquivalente Total (ept_tot),
- Beschäftigungstotal Frauen ($totf$), Beschäftigungstotal Männer ($totm$),
- Vollzeitäquivalente Frauen (ept_f), Vollzeitäquivalente Männer (ept_m).

Die fehlenden Werte werden wie folgt imputiert: Falls die Arbeitsstätte $j \in (S_{trim} \setminus R_{trim}) \cap U_h$, dann

$$y_{trim,j}^* = \begin{cases} evol_h \times y_{trim1,j} & \text{falls } S_{trim1} \text{ existiert} \\ evol_h^0 \times y_{cadre_{t_0},j} & \text{falls } S_{trim1} \text{ nicht existiert und } j \text{ existiert in } T0, h \in \mathcal{H}, \\ y_{cadre_{t_1},j} & \text{sonst} \end{cases} \quad (12)$$

wobei $y_{trim,j}^*$ der imputierte Wert, $y_{trim1,j}$ der Wert (Total oder EPT) aus dem Vorquartal, $y_{cadre_{t_0},j}$ und $y_{cadre_{t_1},j}$ die entsprechenden Werte (Total oder EPT) aus den Rahmen in den Zeitpunkten $T0$ und $T1$ sind. $evol_h^0$ ist die Entwicklung basierend auf der Formel (11), wobei jedoch im Nenner die Vollzeitäquivalente aus dem Initialrahmen $T0$ summiert werden.

Tabelle 11 Notation der Beschäftigten-Variablen aus der Erhebung

Variable	Arbeitszeit in %	Frauen	Männer	Total
Vollzeit	[90-100]	<i>fvz</i>	<i>mvz</i>	<i>totvz</i>
Teilzeit1	[50-90[	<i>ftz1</i>	<i>mtz1</i>	<i>tottz1</i>
Teilzeit2	[15-50[	<i>ftz2</i>	<i>mtz2</i>	<i>tottz2</i>
Teilzeit3	[0-15[	<i>ftz3</i>	<i>mtz3</i>	<i>tottz3</i>

Für die anderen Beschäftigten-Erhebungswerte gemäss Tabelle 11 und nach Einsetzung der ersten Variablen ($tot_{trim,j}^*$, $ept_tot_{trim,j}^*$, $totf_{trim,j}^*$, $totm_{trim,j}^*$, $ept_f_{trim,j}^*$ und $ept_m_{trim,j}^*$), werden die Quotienten für jede dieser Variablen wie folgt berechnet:

$$\begin{aligned}
ratio_{totvz,h} &= \frac{\sum_{k \in R_{job,trim} \cap U_h} totvz_{trim,k}}{\sum_{k \in R_{job,trim} \cap U_h} tot_{trim,k}}, h \in \mathcal{H} \\
ratio_{tottz1,h} &= \frac{\sum_{k \in R_{job,trim} \cap U_h} tottz1_{trim,k}}{\sum_{k \in R_{job,trim} \cap U_h} tot_{trim,k}}, h \in \mathcal{H} \\
ratio_{tottz2,h} &= \frac{\sum_{k \in R_{job,trim} \cap U_h} tottz2_{trim,k}}{\sum_{k \in R_{job,trim} \cap U_h} tot_{trim,k}}, h \in \mathcal{H} \\
ratio_{tottz3,h} &= \frac{\sum_{k \in R_{job,trim} \cap U_h} tottz3_{trim,k}}{\sum_{k \in R_{job,trim} \cap U_h} tot_{trim,k}}, h \in \mathcal{H} \\
&\text{etc...}
\end{aligned} \tag{13}$$

Die eingesetzten Variablen werden wie folgt notiert und berechnet

$$\begin{aligned}
totvz_{trim,j}^* &= ratio_{totvz,h} \times tot_{trim,j}^* \\
tottz1_{trim,j}^* &= ratio_{tottz1,h} \times tot_{trim,j}^* \\
&\text{etc ...} \\
fvz_{trim,j}^* &= ratio_{fvz,h} \times totf_{trim,j}^* \\
ftz1_{trim,j}^* &= ratio_{ftz1,h} \times totf_{trim,j}^* \\
&\text{etc ...} \\
mvz_{trim,j}^* &= ratio_{mvz,h} \times totm_{trim,j}^* \\
mtz1_{trim,j}^* &= ratio_{mtz1,h} \times totm_{trim,j}^* \\
&\text{etc ...}
\end{aligned} \tag{14}$$

Diese Methode sichert die Konsistenz in den Resultaten im Sinn, dass Summen von den eingesetzten Teilsummen wiederum das eingesetzte Total ergeben (wie z.B. $totvz_{trim,j}^* = mvz_{trim,j}^* + fvz_{trim,j}^*$).

Die Prozedur zur Einsetzung der `clusterid` in den Schichten der `TakeAll_R100` ausserhalb der Profilings ist identisch zur oben beschriebenen Prozedur auf Niveau Arbeitsstätte, doch wird alles auf Niveau `clusterid` berechnet.

6.4 Anpassung der Antwortausfälle mit einer Anpassung der Gewichte

Für die Beschäftigtenvariablen kann es totale Antwortausfälle auf Niveau Unternehmen geben, wenn die Ziehungseinheiten die Daten nicht liefern. Totale Antwortausfälle ausserhalb des Profilings oder der Take-ALL-Schicht `TakeAll_R100` werden mittels einer Gewichtungsanpassung behandelt. Dasselbe gilt für Antwortausfälle in Bezug auf die Variablen der offenen Stellen (`item nonresponse`).

Wenn die Variablen der offenen Stellen betrachtet werden, können Antwortausfälle überall vorkommen, d.h. für Arbeitsstätten des Profilings, für Unternehmen ausserhalb des Profilings, wie auch bei `TakeAll_R100`. Die Modellierung der Antwortwahrscheinlichkeiten geschieht hier für alle Einheiten auf Niveau `clusterid`.

6.5 Anpassung der Antwortausfälle für die Beschäftigtenvariablen

Jede Ziehungseinheit (`clusterid`) hat eine Einschlusswahrscheinlichkeit $\pi_{komb,i}$ (Formel 8). Das Ziehungsgewicht kann wie folgt ausgedrückt werden $w_{0,fin,i} = \pi_{komb,i}^{-1}$.

Seien $S_{\text{trim},1} = S_{\text{trim}} \setminus (\text{TakeAll_R100} \cup \text{PROF})$, die Einheiten ausserhalb des Profilings und der Take-All-Schicht. Die Modellierung der Antwortausfälle hat zum Ziel, neue, für die Anpassung der Antwortausfälle korrigierte Gewichte zu berechnen $w_{1,\text{job},\text{fin},i}$. Man erhält diese, indem auf die Initialgewichte $w_{0,\text{fin},i}$ ⁶ ein Korrekturfaktor angewandt wird, der basierend auf der Methode der Segmentierung CHAID (für CHi-squared Automatic Interaction Detection) berechnet wird. Diese Methode erlaubt es mit erklärenden Hilfsvariablen und der Indikatorvariablen (Antwort ja/nein) eine Menge $\mathcal{Q}$ von homogenen Antwortgruppen zu berechnen (response homogeneity groups: rhg). Diese Gruppen sind disjunkt ($\text{rhg}_q \cap \text{rhg}_{q'} = \emptyset, q \neq q', \text{ mit } q, q' \in \{1, \dots, |\mathcal{Q}|\}$) und decken die ganze Bruttostichprobe ab: $\cup_q \text{rhg}_q = S$. Die erklärenden Variablen gemäss Extrapolationsrahmen, die in die Berechnung des Modells einfließen, sind folgende:

- 47 NOGA-BFS50,
- 7 Grossregionen,
- 4 Grössenklassen (Variable «classeTaille»):
 - classeTaille = 1 falls $\text{betot_cid_pop} = 1$,
 - classeTaille = 2 falls $2 \leq \text{betot_cid_pop} \leq 5$,
 - classeTaille = 3 falls $6 \leq \text{betot_cid_pop} \leq 100$,
 - classeTaille = 4 falls $101 \leq \text{betot_cid_pop}$.
- Eine Identifikatorvariable LD:
 - LD = 1, falls der Quotient $\text{ept_cid_pop}/\text{betot_cid_pop} < 0.5$,
 - LD = 0 sonst.

Die Wahl der 4 Grössenklassen beruht auf explorativen Analysen zum Zusammenhang zwischen Unternehmensgrösse und Antwortquote. Die schliesslich verwendeten Gruppen weisen gut unterscheidbare Antwortquoten auf.

$\mathcal{Q}$ wurde mit den Daten des ersten Quartals nach der Ziehung der Stichprobe, also im 2. Quartal 2015 berechnet und wurde nachher im Laufe der Zeit aus Stabilitätsgründen nicht mehr geändert. Aktualisiert bzw. überprüft wurde $\mathcal{Q}$ erst wieder bei der Neuziehung fürs 1. Quartal 2018. Als Input für den CHAID-Algorithmus wurden folgende Parameter definiert:

- die Antwortquote in jeder Zelle (rhg) der Stichprobe $S_{\text{trim},1}$ ist fixiert auf mindestens 50%.
- die Anzahl Einheiten in jeder Zelle (rhg) der Stichprobe $S_{\text{trim},1}$ muss mindestens 100 Einheiten sein.

Formal lässt sich das bezüglich Antwortausfällen korrigierte Gewicht eines Unternehmens i in der Antworthomogenitätsgruppe $q \in \mathcal{Q}$ als

$$w_{1,\text{job},\text{fin},i} = w_{0,\text{fin},i} \cdot \theta_q^{-1} \quad (15)$$

darstellen. Dabei bezeichnet

- θ_q die in der Antworthomogenitätsgruppe q beobachtete Antwortrate (ungewichtet), welche auch als geschätzte Antwortwahrscheinlichkeit betrachtet werden kann und
- $w_{0,\text{job},\text{fin},i}$ das Stichprobengewicht, respektive ein entsprechend (10) angepasstes Stichprobengewicht falls Unternehmen i aus einer Mutation hervorgeht.

⁶Für Unternehmen, welche aus Mutationen hervorgehen, wird das Initialgewicht mit dem Gewicht aus (10) ersetzt.

6.6 Anpassung der Antwortausfälle für die Variablen zur Anzahl offenen Stellen

Die Korrektur für Antwortausfälle bei der Variablen «Anzahl offenen Stellen», besteht aus einer weiteren Anpassung der Gewichte. Diese Anpassung geschieht für die antwortenden Einheiten (`clusterid`), inklusive Profiling und TakeAll_R100, die die Frage der offenen Stellen beantwortet haben $R_{\text{plavac}, \text{trim}} \subset R_{\text{job}, \text{trim}}$.

Die Gewichte $w_{1, \text{plvac}, \text{fin}, i}$ wurden analog zur Antwortausfallskorrektur des Abschnitts 6.5 basierend auf den Gewichten $w_{1, \text{job}, \text{fin}, i}$ berechnet. Die verwendeten erklärenden Variablen, die ins Modell einfließen, sind dieselben wie bei der Methode, die für die Anpassung bei den Beschäftigtenvariablen verwendet wurde, plus zusätzlich die Variable `prof_TA`, die wie folgt definiert ist:

- `prof_TA` = 1, falls die Einheit i zum Profiling und zur Takeall-Schicht gehört,
- `prof_TA` = 2, falls die Einheit i zum Profiling aber nicht zur Takeall-Schicht gehört,
- `prof_TA` = 3, falls die Einheit i nicht zum Profiling aber zur Takeall-Schicht gehört,
- `prof_TA` = 4, falls die Einheit i weder zum Profiling noch zur Takeall-Schicht gehört.

Die minimale Antwortquote für die Berechnung der `rhg` in Bezug auf die Antwortausfallskorrektur für offene Stellen wurde auf 80% festgelegt. Infolge der gewählten Parameter für die minimalen Antwortquoten $0.5 \times 0.8 = 0.4$, stellt man hiermit sicher, dass die totale Gewichtskorrektur im Fall der offenen Stellen auf 2.5 ($=1/0.4$) beschränkt ist.

Die Korrektur der Antwortausfälle wird auf Niveau `clusterid` vorgenommen. Die offenen Stellen des Profilings auf Niveau Arbeitsstätte müssen auf Niveau `clusterid` aggregiert werden. Für eine `clusterid` im Profiling $i \in \text{PROF}$ mit $\mathcal{E}_i$, der Menge aller Arbeitsstätten der `clusterid` i und $R_{\text{plavac}, \text{trim}, \mathcal{E}_i} \subseteq \mathcal{E}_i$ der Menge der antwortenden Arbeitsstätten der `clusterid` i im Quartal `trim`, sind drei Fälle möglich:

1. Alle Arbeitsstätten haben geantwortet: $R_{\text{plavac}, \mathcal{E}_i} = \mathcal{E}_i$, oder
2. keine der Arbeitsstätten hat geantwortet: $R_{\text{plavac}, \mathcal{E}_i} = \emptyset$, oder
3. ein Teil der Arbeitsstätten hat geantwortet: $R_{\text{plavac}, \mathcal{E}_i} \subset \mathcal{E}_i$. Die `clusterid` wird in diesem Fall als «gemischt» bezeichnet.

Die ersten beiden Fälle sind klar wie sie gehandhabt werden. Der letzte Fall kann nicht als komplette Antwort betrachtet werden. Man muss daher die Arbeitsstätten $j \in (\mathcal{E}_i \setminus R_{\text{plavac}, \mathcal{E}_i})$ mit einer Einsetzung behandeln, um eine komplette `clusterid` i zu erhalten. Die offenen Stellen der Arbeitsstätten werden eingesetzt, indem die antwortenden Arbeitsstätten derselben `clusterid` betrachtet werden. Folgende Methode wird angewandt:

Für das erste Quartal nach der Ziehung wird der Anteil $p_{\text{plavac}, 0}$ der Werte, die nicht Null sind, unter den vorhandenen Antworten je `clusterid` berechnet. Dieser Anteil soll auch unter den einzusetzenden Arbeitsstätten beibehalten werden.

Der Anteil von Null offenen Stellen ist dabei sehr hoch. Insgesamt gab es gemäss offiziellen publizierten BESTA-Resultaten auf rund fünf Millionen Beschäftigte im 2015 im Mittel rund 53'000 und im 2018 rund 73'000 offene Stellen. Es wird folgende Methode zur Einsetzung der Anzahl offener Stellen (plavac_{ij}^*) bei den Arbeitsstätten j der gemischten `clusterid` i verwendet:

$$\text{plavac}_{ij}^* = \begin{cases} \text{betot}_j \times z_i & \text{mit Wahrscheinlichkeit } p_{\text{plavac}, 0} \\ 0 & \text{mit Wahrscheinlichkeit } 1 - p_{\text{plavac}, 0} \end{cases} \quad (16)$$

Dabei ist die Variable z_i der Quotient zwischen der Anzahl offener Stellen der antwortenden Arbeitsstätten, die nicht mit Null offenen Stellen antworten, und der Summe des Beschäftigungstotals derselben Arbeitsstätten.

$$z_i = \frac{\sum_{k \in R_{\mathcal{E}_i}, plavac_{ik} > 0} plavac_{ik}}{\sum_{k \in R_{\mathcal{E}_i}, plavac_{ik} > 0} betot_k} \quad (17)$$

wobei $plavac_{ij}$ die erhobene Anzahl offener Stellen der Arbeitsstätten j der `clusterid` i ist.

Die Liste $\mathcal{L}_{trim,plavac}$ der Arbeitsstätten, welche im ersten Quartal eingesetzt wurden, sowie jene der darauffolgenden Quartale, wird gespeichert. Ab dem zweiten Quartal nach der Ziehung, wird bei den fehlenden Arbeitsstätten des Profilings aus gemischten `clusterids` geschaut, ob diese auf der Liste sind und dort mit Null oder mit $betot_j \times z_i$ eingesetzt worden ist.

Falls eine Arbeitsstätte auf der Liste des letzten Quartals $\mathcal{L}_{trim1,plavac}$ ist, erhält sie den Wert Null, wenn sie schon letztes Quartal den Wert Null eingesetzt bekommen hatte, wenn nicht, erhält sie wieder $betot_j \times z_i$, wobei $betot_j$, die Beschäftigungsanzahl im Rahmen zum Zeitpunkt `trim` ist und z_i mit Hilfe der Werte des aktuellen Quartals berechnet wurde. Falls sie nicht auf der Liste $\mathcal{L}_{trim1,plavac}$ ist, wird das Einsetzungsmodell (16) angewandt. Die Bestimmung der «rhg» zur Anpassung der offenen Stellen wird in jedem Quartal vorgenommen.

7 Kalibrierung und Schätzverfahren

Der Abschnitt über die Kalibrierung und Schätzverfahren wird aufgeteilt in die Definition der Kalibrierungsvariablen und –zellen, die angewandte Methode der Kalibrierung mit dem SAS-Makro CALMAR, sowie die Hochrechnung der Zielvariablen inkl. Schätzung ihrer Varianzen.

Im Detail heisst das:

- Definition der Kalibrierungsvariablen und –zellen (Abschnitt 7.1).
- Robustifizierung der extremen Gewichte nach Anpassung für die Antwortausfälle (Abschnitt 7.2).
- Kalibrierung der Beschäftigungs-Gewichte auf die Grenzen wie sie unter 7.1 beschrieben sind (Abschnitt 7.3).
- Schätzung der Zielvariablen in Untersuchungsbereichen inkl. Varianzschätzung für Beschäftigungsvariablen (Abschnitt 7.4).
- Kalibrierung der Gewichte der offenen Stellen auf die Grenzen wie sie unter 7.1 beschrieben sind (Abschnitt 7.5).
- Schätzung der Zielvariablen in Untersuchungsbereichen inkl. Varianzschätzung für die offenen Stellen (Abschnitt 7.6).

7.1 Kalibrierungsvariablen

Die Kalibrierungsgrenzen sind auf Niveau Arbeitsstätten wie folgt definiert

- `cal_bfs50_h_TA_l`, ist eine Kreuzung der Variablen NOGA-BFS50 und der Variable `TakeA11_R100`, mit $h \in \mathcal{H}$ (wie in Tabelle 17 definiert) und $l \in \{0, 1\}$, $l = 1$, falls die Arbeitsstätte zum `TakeA11_R100` gehört.

- $\text{cal_region_g_secteur_k_TA_l}$, ist eine Kreuzung zwischen Regionen, Sektoren und der TakeAll_R100-Variable, mit $g \in G$ (wie in Tabelle 12 definiert), $k \in \{2, 3\}$ und $l \in \{0, 1\}$ mit $k = 2$ resp. 3, falls die Arbeitsstätte zum Sektor 2 resp. 3 gehört.

Die Kalibrierungsgrenzen sind die Summen der Anzahl Beschäftigten (betot) im Rahmen der Arbeitsstätten für jede Zelle der Kalibrierung. Im Fall der Gewichtung für die Beschäftigtenvariablen könnten die TakeAll_R100-Schichten und das Profiling eigentlich ganz aus der Kalibrierung entfernt werden, denn beide Gruppen erhalten ein Hochrechnungsgewicht von 1. Doch entfernt wurde nur das Profiling (siehe Abschnitt 7.3). Die Einheiten der TakeAll_R100 ($l = 1$) Zellen bleiben drin, die Gewichte werden durch die Kalibrierung jedoch nicht verändert. In der Tat sind die Grenzen in der Stichprobe und der Population identisch, da ja die TakeAll_R100 Schichten voll erhoben werden und die Antwortquote nach dem statistischen Datenaufbereitungsprozess 100% beträgt.

Tabelle 12 Liste der Regionen $g \in G = \{1, \dots, 12\}$ für die Kalibrierung

Identifikator (g)	Regionen	Aufstockung
1	VS	nein
2	VD	ja
3	GE	ja
4	BE, FR, SO, JU	nein
5	NE	ja
6	BS, BL, AG	nein
7	ZH (ohne zs)	nein
8	zs	ja
9	GL, SH, AR, AI, GR, TG	nein
10	SG	ja
11	LU, UR, SZ, OW, NW, ZG	nein
12	TI	nein
zs steht für Zürich Stadt ab 2018: Zusätzlich Stadt Winterthur wi, als Nr.13 mit Aufstockung, anstatt in Nr.7		

7.2 Robustifizierung der Gewichte

Die Prozedur der Robustifizierung wird im BESTA-Methodenbericht (OFS: Renaud A., Panchard Ch., Potterat J., 2008, Abschnitt 5.1) beschrieben. Es wurde eine leichte Anpassung vorgenommen, was die Variablen zur Bestimmung der extremen Werte betrifft. Neu werden die Beschäftigten-Variablen « betot » aus dem Extrapolationsrahmen auf Niveau Arbeitsstätten und die Erhebungsvariable « tot » (Beschäftigungstotal aus der Erhebung) verwendet und einander gegenüber gestellt. Dies ist damit zu begründen, dass die Beschäftigungsanzahl auch für die Kalibrierung verwendet wird. Vorher waren es in beiden Fällen (Robustifizierung und Kalibrierung) die Anzahl Vollzeitäquivalente (ept_cid_pop bzw. ept_tot).

7.3 Kalibrierung der Gewichte für die Beschäftigtenvariablen

Die Kalibrierung erfolgt auf die Zellen, wie sie in Abschnitt 7.1 beschrieben sind. Mit der Revision im Jahr 2015 wurden mehrere Tests vorgenommen. In einem ersten Schritt wurde analysiert, inwieweit der zusätzliche Gewichtungsschritt für die Korrektur von Antwortausfällen, gegenüber dem Ansatz «Kalibrierung und Antwortausfallskorrektur in einem Schritt» eine Verbesserung bringt. Hierzu wurden Hochrechnungen von Variablen aus dem Stichprobenrahmen betrachtet. Aus der relativen Differenz zu den effektiven, bekannten Totalen, liess sich eine Biasschätzung berechnen. So wurden die Werte für eine Schätzung der Totale der Variable EPT basierend auf einer Kalibrierung der Variable betot mit bzw. ohne Antwortkorrektur berechnet. Die Antwortkorrektur verringert dabei den geschätzten, relativen Bias von 1.7% auf 0.7% und stützte den

Entscheid, die in den vorangehenden Abschnitten vorgestellte Korrektur für Antwortausfälle in den Gewichtungsprozess zu integrieren.

Weiter stellt sich die Frage nach der Wahl der Kalibrierungsvariablen. Hierzu wurden die Varianten Kalibrierung mittels *betot*, mittels *ept*, mittels *betot* und *ept* in Bezug auf die Varianz der Schätzung der Totale der Erhebungsvariablen *tot* und *ept_tot* verglichen (s. Tabelle 13). Es zeigte sich, dass sich die Varianten nicht wesentlich unterscheiden. Die Variable *ept* beruht in vielen Fällen auf einem Modell (siehe dazu [FSO: Assoulin D., Nedyalkova D., 2017](#)). Demgegenüber reflektieren die *betot* effektiv beobachtete Werte, weshalb diese Variable gewählt wurde. Die Weiterentwicklung der Variante mit beiden Kalibrierungsvariablen hätte zudem eine Anpassung der Robustifizierungsprozedur nötig gemacht, was ein weiterer Grund war, dass diese Variante nicht weiterverfolgt worden ist.

Tabelle 13 Schätzung der Variationskoeffizienten CV für die Schätzung des Totals der Beschäftigtenvariablen ($\hat{Y}_{tot}$ und $\hat{Y}_{ept_tot}$) der Erhebung nach Anpassung der Antwortausfälle und nach Kalibrierung.

Kalibrierung auf	CV von $\hat{Y}_{tot}$ (in %)	CV von $\hat{Y}_{ept_tot}$ (in %)
<i>betot</i>	0.351	0.384
<i>ept</i>	0.377	0.329
<i>betot</i> und <i>ept</i>	0.371	0.358

Die Kalibrierungstotale sind also eine Aggregation der Variablen *betot* im Arbeitsstätten Extrapolationsrahmen *cadre_ast_t1* in verschiedenen Zellen. Es gilt anzumerken, dass die Profilingfälle von der Kalibrierung auf die Beschäftigtenvariablen ausgeschlossen werden und wie die *TakeAll_R100* Schichten ein Hochrechnungsgewicht von 1 erhalten. Dies ist so zulässig, da eventuelle Antwortausfälle in diesen Schichten eingesetzt werden, so dass keine Hochrechnung mehr vorgenommen werden muss.

7.4 Varianzschätzung für die Beschäftigtenvariablen

Die Beschäftigtenvariablen sind in Tabelle 14 beschrieben.

Tabelle 14 Beschäftigtenvariablen

Variable	Arbeitszeit in %	Frauen	Männer	Total
Vollzeit	[90-100]	<i>fvz</i>	<i>mvz</i>	<i>totvz</i>
Teilzeit1	[50-90[	<i>ftz1</i>	<i>mtz1</i>	<i>tottz1</i>
Teilzeit2	[15-50[	<i>ftz2</i>	<i>mtz2</i>	<i>tottz2</i>
Teilzeit3	[0-15[	<i>ftz3</i>	<i>mtz3</i>	<i>tottz3</i>
Vollzeitäquivalente		<i>ept_f</i>	<i>ept_m</i>	<i>ept_tot</i>

Die Schätzung des Totals Y für eine Beschäftigtenvariable der Tabelle 14 kann wie folgt ausgedrückt werden.

$$\hat{Y} = \sum_{j \in R} w_{2,job,betot,fin,j} \times y_j, \quad (18)$$

mit $w_{2,job,betot,fin,j}$ dem Gewicht nach Kalibrierung auf *betot* (mit der gewählten Methode aus

Abschnitt 7.3). Die Schätzung des Variationskoeffizienten $CV_{\hat{Y}}$ der Schätzung des Totals Y , kann wie folgt geschrieben werden

$$\begin{aligned}\widehat{CV}_{\hat{Y}} &= \frac{\sqrt{\widehat{\text{Var}}(\hat{Y})}}{\hat{Y}} \\ &= \frac{\sqrt{\widehat{\text{Var}}\left(\sum_{j \in R \setminus (\text{PROF} \cup \text{TakeAll_R100})} w_{2,\text{job,betot,fin},j} \times y_j\right)}}{\hat{Y}}.\end{aligned}\quad (19)$$

Dabei ist zu bemerken, dass die aktuelle Varianzschätzung die Einsetzungsvarianz nicht berücksichtigt, was eine tendenzielle Unterschätzung der Varianz bedeutet. Weitere Details zur Varianzschätzung gegeben eine Kalibrierung sind im BESTA-Methodenbericht (OFS: [Renaud A., Panchard Ch., Poterat J., 2008](#), Abschnitt 6) beschrieben.

Im Aufruf der SAS-Prozedur `proc surveymeans` wird die Option Schicht «`strata`» nicht mehr gesetzt, da der BESTA ein Poisson-Stichprobenplan zugrunde liegt, womit es a priori keine Schichten mit fixen Stichprobengrößen gibt.

In einigen BFS-Erhebungen wird die Schichtoption auch im Rahmen von Varianzschätzungen für Poissonziehungen verwendet, um auf einfache Art allfällige Poststratifizierungseffekte, welche sich im Rahmen der Gewichtung ergeben, zu berücksichtigen. Aufgrund des flexiblen Ansatzes bei der Vorbereitung des Stichprobenplans und der Gewichtung (keine geplanten Schichten, flexible Wahl von Antwort-Homogenitätsklassen) wäre ein entsprechendes Vorgehen in der BESTA zumindest nicht offensichtlich. Andererseits basieren der Stichprobenplan und die Antwortkorrektur hauptsächlich auf denselben Variablen (NOGA, Anzahl Beschäftigte) wie die Kalibrierung, welche im Rahmen der Varianzschätzung berücksichtigt wird. Darum wurde davon ausgegangen, dass die gewählte Alternative, bei der Varianzschätzung auf Schichtungsvariablen zu verzichten, zu nur leicht höheren Varianzschätzungen führt. Dies wird durch die im Abschnitt 7.7 dargestellten Varianzschätzungen bestätigt, welche im erwarteten Bereich liegen.

Zusätzlich muss in der SAS-Prozedur die Option Klumpenstichprobe «`cluster=clusterid`» gesetzt werden, da die Berechnung der Varianz auf den ersten Stichprobeneinheiten (`clusterid`) erfolgt. Die Option «`_total_`» gibt die relevante Anzahl Klumpen (`clusterid`) im Extrapolationsrahmen an ($U \setminus (\text{PROF} \cup \text{TakeAll_R100})$) und bewirkt, dass im Rahmen der Varianzschätzung eine Endlichkeitskorrektur verwendet wird.

7.5 Kalibrierung der Gewichte für die Variablen zur Anzahl offenen Stellen

Eine der Hauptvariablen bei den offenen Stellen ist die Variable der Anzahl offener Stellen «`nbrplvac`». Nach Anpassung für die Antwortausfälle, werden die erhaltenen Gewichte $w_{1,\text{plvac,fin},i}$ auf die Variablen aus Abschnitt 7.1 kalibriert. Es gibt im Gegensatz zu den Beschäftigtenvariablen (Abschnitt 7.2) keine Behandlung von extremen Werten für die Gewichte $w_{1,\text{plvac,fin},i}$ und keine Einsetzung bei fehlenden Werten in der Schicht der `TakeAll_R100` bei den offenen Stellen.

Wegen letzterem Punkt wird der Kalibrierungsrahmen für die offenen Stellen erweitert und besteht aus allen Arbeitsstätten inklusive Profiling und `TakeAll_R100`. Die Kalibrierung wird auf dieselbe Art vorgenommen wie bei den Beschäftigtenvariablen.

7.6 Varianzschätzung für die Variablen zur Anzahl offenen Stellen

Die Abbildung 2 zeigt einen Auszug des Fragebogens und enthält alle aufgeführten Variablen, die bei den offenen Stellen geschätzt werden.

Abbildung 2 Variablen offenen Stellen

2. Anzahl offener Stellen am Ende des Berichtsquartals				Keine offenen Stellen ☐	
3. Personalrekrutierung im Berichtsquartal					
Hatten Sie im Berichtsquartal Schwierigkeiten bei der Rekrutierung von Personal mit: (nur eine Antwort pro Zeile ankreuzen)		Diese Personalkategorie wurde...			
		...leicht gefunden	...schwer gefunden	...nicht gefunden	nicht gesucht/ Suchprozess noch nicht abgeschlossen/ weiss nicht
– Hochschulabschluss (Universität, ETH, Fachhochschule oder gleichwertige Ausbildung)	☐	☐	☐	☐	☐
– höhere Berufsbildung (Meisterprüfung, eidg. Fachausweis, höhere Fachschule usw.)	☐	☐	☐	☐	☐
– Berufslehre oder gleichwertige Ausbildung (eidg. Fähigkeitsausweis, Maturität usw.)	☐	☐	☐	☐	☐
– Obligatorische Schulbildung (ohne nachobligatorische Ausbildung)	☐	☐	☐	☐	☐
4. Voraussichtliche Anzahl Beschäftigte im nächsten Quartal (nur eine Antwort ankreuzen)					
Beibehaltung des Personalbestands ☐		Erhöhung des Personalbestands ☐		Reduktion des Personalbestands ☐	
Im Falle einer Erhöhung oder Reduktion bitte die Veränderung in Anzahl Beschäftigten einschätzen					

Es gibt also eine Unterscheidung von quantitativen und qualitativen Variablen.

- Quantitative Variablen
 - Anzahl offenen Stellen (nbrplvac).
 - Schätzung der Veränderung des Personalbestandes (estimate_nb).
- Qualitative Variablen
 - keinplvac =1, falls keine offenen Stellen, 0 sonst.
 - replvac =1, falls im aktuellen Quartal die Antwort auf die Frage der offenen Stellen vorliegt, 0 sonst (Antwortindikator).
 - Personalrekrutierung: Personal mit Hochschulabschluss (academic_cd), höhere Berufsbildung (p_s_cd), Berufslehre (a_s_cd) und obligatorische Schulbildung (o_s_cd): 1= gefunden ohne Schwierigkeiten / 2= Gefunden mit Schwierigkeiten / 3= Nicht gefunden / 4= Nicht gesucht / 9= keine Antwort.
 - Entwicklung des Personalbestandes (forecasts_cd): 1= Beibehaltung / 2= Erhöhung / 3= Senkung.

Die qualitativen Variablen können durch die Bildung von Indikatorvariablen für jede Modalität der Fragen in quantitative Variablen umgewandelt werden.

Die Schätzung des Totals Y für die quantitativen oder qualitativen Variablen ist

$$\hat{Y} = \sum_{j \in R} w_{2,plvac,betot,fin,j} \times y_j, \quad (20)$$

mit $w_{2,plvac,betot,fin,j}$ dem Gewicht nach Kalibrierung auf die Beschäftigtenvariablen.

Die Schätzung des Variationskoeffizienten $CV_{\hat{Y}}$ des Totalschätzers Y wird wie folgt geschrieben.

$$\widehat{CV}_{\hat{Y}} = \frac{\sqrt{\widehat{\text{Var}}(\hat{Y})}}{\hat{Y}} \quad (21)$$

Die Varianz $\widehat{\text{Var}}(\hat{Y})$ in der Gleichung (21) wird auf der Menge R_{plavac} berechnet, also von allen Einheiten inklusive Profiling, welche auf die Frage der offenen Stellen geantwortet haben.

Auch hier kann der Methodenbericht ([OFS: Renaud A., Panchard Ch., Potterat J., 2008](#), Abschnitt 6) für weitere Details zur Varianzschätzung bei Kalibrierung in der BESTA konsultiert werden. Die Varianzschätzung erfolgt analog zum Fall der Beschäftigtenvariablen mit der SAS-Prozedur `proc surveameans`, mit dem Unterschied, dass für die Variablen der offenen Stellen infolge Antwortausfällen auch die Unternehmen aus dem Profiling und `TakeAll_R100` in die Schätzung einfließen.

7.7 Resultate

Die Tabelle 15 zeigt zur Veranschaulichung der Methoden die provisorischen⁷ Resultate inklusive Genauigkeit für das erste Quartal nach der Erneuerung der Stichprobe, d.h. fürs 2. Quartal 2015, sowie das erste Quartal 2018 für diverse Untersuchungsbereiche (Domains). Die Spalten `EPT_TOT` bzw. `TOT` enthalten die geschätzte Anzahl Vollzeitäquivalente bzw. Anzahl Beschäftigte auf Niveau Arbeitsstätten. Die Spalten `CV_EPT` bzw. `CV_TOT` sind die entsprechenden Variationskoeffizienten (in %).

Es ist ersichtlich, dass die angestrebten Zielgenauigkeiten hinsichtlich der Schätzung von `TOT` fast überall erreicht werden. In beiden Jahren gibt es je zwei Schätzungen mit einem leicht zu hohen `CV_TOT` über 4% (grau hervorgehoben). Für die Schätzung der `EPT_TOT` wurden keine Zielgenauigkeiten definiert. Aber auch dort sind die `CV_EPT` mit Ausnahme von je zwei Fällen (über 4%) sehr präzise. Bei den Genauigkeiten für die Regionen mit einer Erhöhung wird der erwünschte CV von 3% für den Sektor überall eingehalten.

⁷Die offiziellen, bereinigten Resultate sind auf der Webseite des BFS zu entnehmen und können von diesen hier abweichen.

Tabelle 15 Resultate der Erhebung im 1. Quartal nach der Ziehung (2018-1 bzw. 2015-2)

Resultate BESTA, 2018-1					Resultate BESTA, 2015-2							
Domain	Npop	nBrut	nNet	nNetUNT	EPT_TOT	CV_EPT	TOT	CV_TOT	EPT_TOT	CV_EPT	TOT	CV_TOT
TOTAL	626'155	66'164	62'592	20'032	3'894'136.6	0.36	5'005'086.2	0.29	3'843'468.1	0.38	4'873'374.5	0.33
Sektor												
2	95'085	8'306	7'269	5'109	981'634.3	0.44	1'074'226.6	0.40	1'001'378.6	0.49	1'093'081.5	0.45
3	531'070	57'858	55'323	14'923	2'912'502.3	0.46	3'930'859.6	0.35	2'842'089.6	0.49	3'780'293.0	0.40
NOGA												
5.9	354	134	115	67	4'476.3	1.63	4'875.9	1.85	4'499.9	2.50	4'834.8	2.35
10.2	5'062	649	553	240	74'547.3	1.65	86'558.5	1.27	75'414.2	1.51	87'546.1	1.46
13.5	3'062	268	218	197	12'394.5	2.48	14'996.2	1.87	13'580.0	2.67	16'201.6	2.68
16.8	9'921	634	544	499	59'976.0	1.43	68'401.6	1.39	62'137.0	1.56	70'579.6	1.51
19.2	806	189	165	126	28'092.4	1.43	29'915.9	1.16	28'332.5	1.79	30'394.9	1.60
21	286	84	78	53	44'076.7	1.65	48'560.3	1.46	41'755.7	1.28	44'016.9	1.28
22.3	2'310	438	398	303	37'031.5	1.27	40'015.2	1.08	38'161.0	1.54	40'888.8	1.39
24.5	7'905	588	502	450	88'863.5	1.25	96'383.2	1.06	93'623.8	1.29	101'539.7	1.24
26	2'220	460	410	275	100'381.6	1.30	108'476.8	1.28	103'119.5	1.33	109'222.9	1.25
27	891	187	164	121	30'318.8	1.11	31'852.3	1.12	29'447.0	1.81	30'141.5	1.79
28	2'200	385	296	255	72'940.0	1.35	77'952.6	1.36	75'872.4	1.15	80'899.4	1.08
29.3	476	83	73	56	15'802.6	1.37	16'274.9	1.49	15'048.7	3.19	15'718.5	3.16
31.3	7'614	578	507	405	49'086.6	1.77	55'746.4	1.31	49'451.7	1.50	55'705.3	1.18
35	1'114	421	408	206	25'073.2	1.40	28'505.6	1.11	25'933.7	1.03	29'579.3	1.04
36.9	2'048	358	324	204	16'444.6	1.95	19'782.8	1.97	15'939.9	2.33	18'154.2	1.86
41.2	9'558	1'001	894	541	106'056.9	1.36	113'007.5	1.27	109'958.0	1.34	116'943.7	1.23
43	39'258	1'849	1'620	1'111	216'271.5	1.36	236'920.9	1.19	214'950.9	1.67	234'714.2	1.55
45	17'106	996	904	533	80'049.1	1.59	91'232.6	1.34	80'558.4	1.75	90'457.4	1.46
46	27'985	1'830	1'642	940	199'435.2	1.33	228'465.0	1.14	207'407.0	1.49	239'266.7	1.43
47	51'721	9'860	9'622	1'320	231'916.7	1.70	309'020.2	1.20	242'552.8	2.07	323'242.7	2.00
49	11'961	1'243	1'157	295	109'224.6	1.34	124'076.0	1.20	104'821.1	2.27	119'981.3	1.92
50.1	430	101	88	66	15'354.2	4.29	17'231.2	4.20	13'922.1	2.98	15'841.8	3.06
52	2'452	707	678	291	52'499.9	1.01	59'254.1	0.80	49'566.9	1.66	56'246.9	1.45
53	3'083	2'597	2'593	34	30'408.0	2.94	51'719.9	7.23	31'585.6	3.10	47'582.6	4.70
55	5'846	695	624	401	63'857.1	1.95	78'140.2	1.83	63'458.5	2.36	77'030.0	2.25
56	25'376	2'606	2'427	656	123'359.9	4.81	180'967.8	2.95	120'084.2	2.48	172'340.8	2.23
58.6	5'056	368	332	188	25'069.8	2.76	34'948.1	2.01	24'429.2	2.67	34'292.2	2.30
61	1'206	743	737	44	26'459.7	0.74	28'022.2	1.07	27'568.4	2.21	29'143.2	2.04
62.3	19'078	845	764	591	92'555.9	2.35	109'285.4	2.12	85'388.1	2.53	98'028.4	2.23
64	8'989	2'535	2'497	497	105'449.2	0.84	117'666.5	0.70	118'444.5	0.96	130'025.0	0.71
65	1'289	769	750	92	43'487.5	1.99	49'969.1	1.54	42'975.5	1.05	48'515.4	0.78
66	11'075	1'315	1'209	490	56'534.5	2.01	65'729.2	1.61	53'546.0	2.19	62'510.6	1.85
68	17'757	717	624	426	40'072.3	3.75	67'605.2	2.60	44'884.6	5.24	57'336.4	3.88
69	23'506	731	638	532	65'699.5	2.40	87'792.6	2.08	62'388.3	2.75	82'108.3	2.36
70	23'632	1'002	923	490	97'466.9	2.44	121'738.3	1.50	91'451.2	2.53	111'479.9	2.27
71	26'643	1'098	1'013	789	111'079.8	1.85	134'756.0	1.42	104'674.3	1.82	124'584.6	1.54
72	1'854	238	212	135	22'486.4	2.12	25'476.9	2.04	20'827.4	2.64	23'374.1	2.44
73.5	25'853	664	546	435	46'970.2	3.00	66'958.0	1.74	44'105.6	2.62	63'924.3	2.48
76	3'577	1'469	1'400	278	84'526.8	3.60	112'741.7	2.67	86'886.7	5.81	116'643.9	4.55
79.2	25'444	2'515	2'370	741	137'623.6	2.50	216'845.8	1.70	137'781.3	2.43	212'074.6	1.64
84	8'654	5'014	5'002	339	161'728.9	0.69	200'136.8	0.59	159'646.6	1.03	200'120.7	0.95
85	32'567	8'892	8'791	537	226'911.2	2.13	357'940.1	1.35	213'019.4	2.00	335'396.9	1.46
86	57'418	1'592	1'492	784	287'993.0	1.67	411'004.8	1.07	264'103.1	1.78	374'962.6	1.42
87	3'894	1'071	1'053	546	144'369.6	1.23	196'315.0	0.93	131'431.5	1.69	181'138.6	1.28
88	11'319	2'282	2'213	1'006	75'761.1	2.08	120'648.2	1.20	67'363.3	2.69	107'870.4	1.98
90.3	24'835	1'455	1'282	608	56'609.7	3.88	99'930.9	2.32	51'852.0	3.71	92'086.0	2.76
94.6	51'464	1'908	1'740	839	97'841.9	3.11	166'442.0	2.32	96'058.9	3.33	152'213.8	2.10
Grossregion												
12'054	14'321	13'473	4'219		756'868.3	0.72	940'459.3	0.61	112'529	13'889	12'828	3'834
2	122'097	14'392	13'690	4'018	781'342.5	0.77	1'023'764.4	0.55	117'139	14'431	13'581	3'801
3	79'251	7'794	7'452	2'057	520'263.7	1.15	676'419.9	1.06	74'664	7'815	7'256	1'953
4	114'302	12'242	11'610	3'761	778'411.8	0.72	1'004'909.9	0.58	108'659	11'309	10'539	3'190
5	85'913	8'948	8'389	3'306	498'047.3	1.02	640'662.2	0.73	82'601	9'178	8'474	3'440
6	68'467	5'685	5'369	1'772	373'771.5	1.41	499'396.3	0.99	63'861	5'570	5'081	1'619
7	36'071	2'782	2'609	899	185'631.5	1.56	223'474.2	1.18	32'193	2'884	2'655	840
Sektor X Gr												
2-1	16'139	1'586	1'310	934	146'693.6	1.03	157'877.6	0.95	164'889	1'580	1'359	949
2-2	21'884	1'933	1'691	1'187	236'933.5	1.01	261'234.6	0.94	218'669	1'962	1'713	1'201
2-3	11'843	843	756	503	161'547.1	1.25	173'500.3	1.13	119'566	803	722	461
2-4	13'353	1'097	956	679	133'107.9	1.66	147'070.4	1.56	130'564	1'266	1'105	818
2-5	16'078	1'645	1'440	1'093	169'635.4	0.95	186'554.8	0.86	160'885	1'608	1'421	1'056
2-6	10'132	716	616	403	103'852.2	1.68	114'293.8	1.64	103'115	743	643	415
2-7	5'008	346	301	210	49'608.9	1.95	52'550.1	1.75	53'171	344	306	209
3-1	96'390	12'303	11'518	2'900	589'801.3	0.89	747'643.4	0.82	103'565	12'741	12'114	3'270
3-2	95'255	12'498	11'890	2'614	543'144.8	1.27	748'000.0	0.97	100'228	12'430	11'977	2'817
3-3	62'821	6'972	6'500	1'450	364'689.0	1.47	488'037.0	1.20	67'295	6'991	6'730	1'596
3-4	95'306	10'212	9'583	2'511	629'756.5	0.91	833'014.9	0.71	101'248	10'976	10'505	2'943
3-5	66'523	7'533	7'034	2'347	317'079.9	1.17	434'437.6	0.99	69'828	7'340	6'968	2'250
3-6	53'729	4'854	4'465	1'216	259'142.1	1.82	359'842.8	1.39	58'152	4'942	4'726	1'357
3-7	27'185	2'538	2'354	630	138'476.0	2.75	169'517.4	2.64	30'754	2'438	2'303	690
KantErh												
11 (VD)	55'229	6'882	6'517	1'899	344'225.6	1.10	429'709.4	0.98	52'027	6'716	6'197	1'757
12 (GE)	39'335	4'749	4'395	1'618	283'269.4	1.12	342'715.5	0.80	36'606	4'573	4'173	1'471
21 (NE)	12'567	2'210	2'068	774	82'788.1	1.28	102'149.5	0.94	11'980	2'373	2'210	897
41 (zs)	43'847	5'077	4'833	1'580	363'383.8	1.10	467'162.6	0.93	41'057	4'861	4'516	1'429
42 (w)	7'676	1'391	1'270	484	54'879.5	1.59	72'194.8	1.08				
51 (SG)	34'620	4'265	3'931	1'738	225'000.8	1.04	289'609.1	0.82	33'235	4'484	4'101	1'921
Sektor X KantErh												
2-11	7'819	804	700	506	65'313.3	1.70	70'498.5	1.60	7'799	865	704	522
2-12	4'357	488	400	293	44'049.8	1.20	46'221.1	1.16	4'196	450	363	275
2-21	2'419	387	343	264	32'078.5	1.42	33'940.0	1.32	2'419	470	403	339
2-41	2'569	383	328	264	26'433.0	2.42	28'473.6	2.46	2'623	345	293	229
2-42	865	211	175	139	11'249.8	1.81	12'086.6	1.59				
2-51	6'717	905	800	623	80'120.3	1.00	88'051.2	0.85	6'740	937	803	635
3-11	47'410	6'078	5'817	1'393	278'912.3	1.25	359'210.9	1.10	44'228	5'851	5'493	1'235
3-12	3											

8 Schlussfolgerungen

Mit der Revision 2015 sollte bei der Definition der BESTA-Grundgesamtheit erstmals die Umstellung bzw. die Erweiterung des Unternehmens- und Beschäftigtenuniversums gemäss der STATENT berücksichtigt werden. Zudem wurde die BESTA in das Stichprobenverwaltungssystem des BFS integriert. Dies erlaubt die Koordination der BESTA-Stichproben im Laufe der Zeit und die Koordination der BESTA-Stichproben mit den Stichproben anderer Unternehmenserhebungen, welche ebenfalls im SVS gezogen werden.

Im SVS entsprechen die Ziehungseinheiten des privaten Sektors den Unternehmen, im öffentlichen Sektor den Verwaltungseinheiten. Um die BESTA ins SVS integrieren zu können, wurde sie als 2-stufige Erhebung konzipiert: In einer ersten Stufe werden Unternehmen und Verwaltungseinheiten, vereinfachend oft einfach als Unternehmen bezeichnet, gezogen. Zu jedem im SVS gezogenen Unternehmen werden in der zweiten Stufe auch sämtliche Arbeitsstätten in die Stichprobe aufgenommen (Klumpenstichprobe). Die Schätzungen erfolgen weiterhin auf Niveau Arbeitsstätten.

Das Profiling erhebt für Unternehmen, welche Bestandteil des Profilings sind, quartalsweise Beschäftigteninformationen. Aus Effizienzgründen wird die Profiling-Information in die quartalsweisen BESTA-Schätzungen integriert. Dies führt zu einer Reduktion des Stichprobenumfangs der restlichen Unternehmen.

Obwohl die Stichprobe im SVS als Poissonstichprobe gezogen wird, erlauben die Schichtungs- und Allokationstechniken für einfache geschichtete Zufallsstichproben, die Berechnung von Ziehungswahrscheinlichkeiten auf Ebene Unternehmen. Dabei wurde im Falle des nationalen Planes

- auf Ebene Unternehmen
- mit einer Schichtung NOGA-BFS50 gekreuzt mit Grössenklassen
- mit einer Zielvorgabe von $CV = 4\%$ für eine Genauigkeit der Schätzung der Anzahl Beschäftigten auf Niveau NOGA-BFS50

gearbeitet. Dies führte 2018 auf eine erwartete Brutto-Stichprobengrösse von 16'043 Unternehmen.

Da die Schätzungen der BESTA auf Ebene Arbeitsstätte erfolgen, war es wichtig, die entsprechend dem Stichprobenplan resultierenden Präzisionen für Schätzungen auf Ebene Arbeitsstätte zu untersuchen. Dabei zeigte es sich, dass sich die zu erwartenden Variationskoeffizienten der auf Niveau Arbeitsstätte erzielten Schätzungen für die Anzahl Beschäftigten nach Branchen NOGA-BFS50 nicht wesentlich von jenen der Schätzungen auf Niveau Unternehmen unterscheiden. Die Unterschiede kommen von den Mehrbetriebsunternehmen, welche Arbeitsstätten aus unterschiedlichen NOGA-BFS50 Branchen haben.

Analog zum nationalen Plan wurden für Regionen, welche ihre Stichprobe erhöhen, regionale Stichprobenpläne erstellt. Die effektive Ziehungswahrscheinlichkeit wurde für jedes Unternehmen mittels dem Maximum der Ziehungswahrscheinlichkeiten des nationalen und der regionalen Pläne berechnet. Auf diese Weise wird sichergestellt, dass der resultierende Stichprobenplan sowohl nationale wie regionale Präzisionsziele erfüllt. Zudem kann der Effekt der regionalen Erhöhungen auf die erwartete Stichprobengrösse recht einfach und transparent berech-

net werden. So ergab sich 2018 für den kombinierten Stichprobenplan eine erwartete Brutto-Stichprobengrösse von 18'730 Unternehmen, das heisst, die regionalen Erhöhungen bewirken eine erwartete Vergrösserung der Stichprobe um 2687 Unternehmen.

Einzelne sehr grosse Unternehmen können einen grossen Einfluss auf die Resultate haben. Aus Effizienzgründen ist es vorteilhaft, sowohl Stichprobenfehler wie auch Fehler durch Antwortausfälle bei solchen Unternehmen zu vermeiden. Aus diesem Grund wurde im Falle des nationalen Stichprobenplanes mittels einem iterativen Algorithmus eine Menge von ca. 500 Unternehmen identifiziert, welche einerseits a priori voll erhoben werden müssen (kein Stichprobenfehler) und für welche andererseits Antwortausfälle unbedingt zu vermeiden sind, damit das Erreichen der Präzisionsziele des nationalen Planes gewährleistet werden kann. Analog zum nationalen Plan wurde dieses Vorgehen auch für die regionalen Pläne angewandt.

Zwischen den einzelnen Erhebungen werden im Betriebs- und Unternehmensregister BUR laufend Aktualisierungen vorgenommen. Diese betreffen einerseits den Bestand (Mutationen wie Schliessungen, Neugründungen, Fusionen, Splits) von Unternehmen und Arbeitsstätten, andererseits auch die Variablen (Änderung der NOGA, Adressen, Kanton, Beschäftigte etc.). In der BESTA wird die Anpassung der Hochrechnung und der Stichprobe an solche Änderungen im Rahmen des Reportings geregelt. Dabei ist es wichtig, zuverlässige quartalsweise Evolutions-schätzungen produzieren zu können, welche nicht durch administrative Prozesse beeinflusst werden. Aus diesem Grunde werden Populationsveränderungen im Register bei Nicht-Profiling Einheiten im Rahmen der Schätzung weitgehend ausgeblendet und nur für die Anpassung der Stichprobe berücksichtigt. Die Methode der Gewichtsteilung ist dabei zentral bei der Hochrechnung von Stichprobeneinheiten, welche aus Mutationen hervorgehen. Demografische Veränderungen in der Population der Profiling-Unternehmen werden generell direkt in die aktuellsten Quartalsschätzungen der BESTA übernommen.

Die im Rahmen der Hochrechnungsarbeiten durchgeführten Analysen hinsichtlich der Hochrechnung der Anzahl Beschäftigten haben darauf hingewiesen, dass es zur Minderung des Biasrisikos vorteilhaft ist, vor der Kalibrierung eine Korrektur für Antwortausfälle einzuführen. Im Falle von Profilingunternehmen und Unternehmen, für welche im Rahmen der Planausarbeitung eine 100%-Antwortwahrscheinlichkeit vorausgesetzt wurde, werden Antwortausfälle mittels Einsetzungen behandelt. Dabei steht die Stabilität der Resultate im zeitlichen Verlauf im Vordergrund. Für die übrigen Unternehmen erfolgte eine Gewichtskorrektur mittels Antwortraten in Antworthomogenitätsgruppen auf Niveau Unternehmen. Die Antworthomogenitätsgruppen wurden mittels einem Segmentierungsalgorithmus (CHAID) basierend auf Unternehmensgrösse, NOGA-BFS50, Grossregion und mittlerem Beschäftigungsgrad festgelegt. Die Antwortenden bezüglich der Variablen zu den offenen Stellen bilden eine Untergruppe der Antwortenden zu den Beschäftigtenvariablen (Nettostichprobe). Aus diesem Grund wurde für die Hochrechnung der Variablen zu den offenen Stellen eine bedingte Antwortwahrscheinlichkeit (Antwort offene Stellen, gegeben das Unternehmen hat bezüglich der Beschäftigung geantwortet) geschätzt. Diese Schätzung wird für eine weitere Gewichtskorrektur verwendet. Daraus ergibt sich, dass in der BESTA einerseits ein Gewicht für die Hochrechnung von Beschäftigtenvariablen und andererseits eines für die Hochrechnung von Variablen zu offenen Stellen zu Verfügung gestellt wird.

Sowohl die Gewichte für die Hochrechnung von Beschäftigtenvariablen wie auch jene für die Hochrechnung von Variablen zu offenen Stellen werden mittels einer Kalibrierung angepasst. Da die Auswertungen auf Niveau Arbeitsstätte erfolgen, werden auch die Gewichte auf Niveau Arbeitsstätten kalibriert. Die Kalibrierung erfolgt mittels der Hilfsvariablen *betot* (Anzahl Beschäftigte) und berücksichtigt die wichtigsten Hochrechnungsbereiche (z.B. NOGA-BFS50). Die

Variable *betot* aus dem Rahmen korreliert erfahrungsgemäss gut mit der entsprechenden Variablen aus der Erhebung, womit durch die Kalibrierung für diese Variable mit einer wesentlichen Varianzreduktion gerechnet werden darf. Für andere Variablen, wie die offenen Stellen, ist dies teilweise leider nicht der Fall.

Die Varianzschätzung erfolgt mittels der SAS Prozedur *Surveymeans*. Sie wurde im Rahmen der Revision nur minimal angepasst und berücksichtigt, dass es sich bei der BESTA um eine Klumpenstichprobe handelt. Um dem Effekt der Kalibrierung auf die Varianz Rechnung zu tragen, werden die Zielvariablen für die Varianzschätzung durch die entsprechenden Kalibrierungsresiduen ersetzt. Die für die Erhebung des ersten Quartals 2018 erhaltenen Resultate weisen darauf hin, dass die geplanten Präzisionen (Variationskoeffizienten) sowohl national wie auch regional im Grossen und Ganzen eingehalten oder gar übertroffen werden.

Ausblick:

Im Rahmen der Revision konzentrierten sich die methodischen Arbeiten in erster Linie auf die Ausarbeitung des Stichprobenplans sowie die Korrektur für die Antwortausfälle. Andere Aspekte wie die Kalibrierung, Einsetzungsmethodik oder die Varianzschätzung wurden im Rahmen der Revision – auch aus Zeitgründen – nur minimal überarbeitet und an die neuen Gegebenheiten angepasst. Es gäbe hier aber sicherlich noch Spielraum für Weiterentwicklungen. Beispielsweise wäre es erstrebenswert, die mit den Einsetzungen verbundene Varianz in der Varianzschätzung berücksichtigen zu können, was in der aktuellen Varianzschätzung noch nicht der Fall ist.

Anhang

A Beschrieb Variablen

A.1 Definitionen Betriebstypen

Tabelle 16 Betriebstypen

betyp	BESTA_ relevant	Beschrieb	betyp	BESTA_ relevant	Beschrieb
E13	1	Privatsektor mit Beschäftigten	E00	0	Administrative Verbindungseinheit
E20	1	öffentlicher Sektor: Bund	E03	0	Administrative Einheit des öffentlichen Sektors
E21	1	öffentlicher Sektor: Kanton	E05	0	Administrative Einheit MWST
E22	1	öffentlicher Sektor: Bezirk	E06	0	Administrative Einheit generell
E23	1	öffentlicher Sektor: Gemeinde	E07	0	Administrative Einheit AHV
E24	1	öffentlich-rechtliche Körperschaft	E12	0	Privatsektor ohne Beschäftigten und mit Aktivität
E27	1	öffentliches Unternehmen mit Beschäftigten	E15	0	Privatsektor ohne Beschäftigte und Aktivität
E29	1	Unternehmenseinheit des öffentlichen Sektors	E30	0	Unternehmenseinheit einem ausländischen Staat angehörend
E01	0	Agrarsektor mit Beschäftigten	E50	0	in Liquidation / Konkurs
E02	0	Agrarsektor mit Beschäftigten unter oblig. Schwelle	E51	0	Inaktiv
E14	0	Privatsektor mit Beschäftigten unter oblig. Schwelle	E91	0	Löschung
E17	0	Reaktivierung	E99	0	Löschungen, weil doppelt
E18	0	Neuaufnahme			
E28	0	öffentlicher Sektor ohne Beschäftigte			

A.2 Definitionen NOGA

Tabelle 17 Liste der NOGA Gruppen NOGA-BFS50

NOGA Gruppen NOGA-BFS50	Bedeutung
Sektor 1	
01 - 03	Landwirtschaft, Forstwirtschaft und Fischerei ¹
Sektor 2	
05 - 09	Bergbau und Gewinnung von Steinen und Erden
10 - 12	Herstellung von Nahrungsmitteln und Tabakerzeugnissen
13 - 15	Herstellung von Textilien und Bekleidung
16 - 18	Herstellung von Holzwaren, Papier und Druckerzeugnissen
19 + 20	Kokerei, Mineralölverarbeitung und Herstellung von chemischen Erzeugnissen
21	Herstellung von pharmazeutischen Erzeugnissen
22+23	Herstellung von Gummi- und Kunststoffwaren
24+25	Herstellung von Metallerzeugnissen
26	Herstellung von Datenverarbeitungsgeräten und Uhren
27	Herstellung von elektrischen Ausrüstungen
28	Maschinenbau
29+30	Fahrzeugbau
31 - 33	Sonstige Herstellung von Waren, Reparatur und Installation
35	Energieversorgung
36 - 39	Wasserversorgung, Beseitigung von Umweltverschmutzungen
41 - 42	Hoch- und Tiefbau
43	Sonstiges Ausbaugewerbe
Sektor 3	
45	Handel und Reparatur von Motorfahrzeugen
46	Grosshandel
47	Detailhandel
49	Landverkehr und Transport in Rohrfernleitungen
50 - 51	Schifffahrt und Luftfahrt
52	Lagerei sowie Erbringung von sonstigen Dienstleistungen für den Verkehr
53	Post-, Kurier- und Expressdienste
55	Beherbergung
56	Gastronomie
58 - 60	Verlagswesen, audiovisuelle Medien und Rundfunk
61	Telekommunikation
62+63	Informationstechnologie und Informationsdienstleistungen
64	Erbringung von Finanzdienstleistungen
65	Versicherungen
66	Mit Finanz- und Versicherungsdienstleistungen verbundene Tätigkeiten
68	Grundstücks- und Wohnungswesen
69	Rechts- und Steuerberatung, Wirtschaftsprüfung
70	Unternehmensverwaltung und -führung; Unternehmensberatung
71	Architektur- und Ingenieurbüros
72	Forschung und Entwicklung
73 - 75	Sonstige freiberufliche, wissenschaftliche und technische Tätigkeiten
77 + 79 - 82	Erbringung von sonstigen wirtschaftlichen Dienstleistungen
78	Vermittlung und Überlassung von Arbeitskräften
84	Öffentliche Verwaltung
85	Erziehung und Unterricht
86	Gesundheitswesen
87	Heime (ohne Erholungs- und Ferienheime)
88	Sozialwesen (ohne Heime)
90 - 93	Kunst, Unterhaltung und Erholung
94 - 96	Erbringung von sonstigen Dienstleistungen
97+ 98	Private Haushalte als Arbeitgeber und Hersteller von Waren ¹
99	Exterritoriale Organisationen ¹

¹ kursiv: nicht relevant für BESTA

Literatur

- [Baillargeon und Rivest 2009] BAILLARGEON, Sophie ; RIVEST, Louis-Paul: A General Algorithm for Univariate Stratification. In: *International Statistical Review* 77,3 (2009), S. 331–344
- [Baillargeon und Rivest 2011] BAILLARGEON, Sophie ; RIVEST, Louis-Paul: The construction of stratified designs in R with the package stratification. In: *Survey Methodology, Statistics Canada, Catalogue No. 12-001-X* Vol. 37, No. 1 (2011), S. 53–65
- [Deville und Särndal 1992] DEVILLE, J.-C. ; SÄRNDAL, C.-E.: Calibration estimators in survey sampling. In: *Journal of the American Statistical Association* 87 (1992), S. 376–382
- [FSO: Assoulin D., Nedyalkova D. 2017] FSO: ASSOULIN D., NEDYALKOVA D.: Construction of full-time equivalents for the Swiss structural business statistics 2011. FSO number: 338-0077 / Federal Statistical Office. Neuchâtel, 2017. – FSO methodology reports. – URL <https://www.bfs.admin.ch/bfs/en/home/services/recherche/methodological-reports.assetdetail.2963860.html>
- [Hidiroglou 1986] HIDIROGLOU, M. A.: The Construction of a Self-Representing Stratum of Large Units in Survey Design. In: *The American Statistician* Vol. 40, no. 1 (1986), S. 27–31
- [J.-C.Deville und Lavallée 2006] J.-C.DEVILLE ; LAVALLÉE, P.: Indirect Sampling: The Foundations of the Generalized Weight Share Method. In: *Survey Methodology* Vol. 32, no. 2 (2006), S. 165–176
- [Luzi O. 2007] LUZI O., et a.: *Recommended Practices for Editing and Imputation in Cross-Sectional Business Survey*. EDIMBUS, 2007
- [OFS: Graf M. 2001] OFS: GRAF M.: Désaisonnalisation. Aspects méthodologiques et application à la statistique de l'emploi / Office Fédéral de la Statistique. Neuchâtel, 2001. – BFS Methodenbericht. – URL <https://www.bfs.admin.ch/bfs/fr/home/services/recherche/rapports-methodologiques.assetdetail.341905.html>
- [OFS: Renaud A. 2008] OFS: RENAUD A.: Statistique de l'emploi. Révision 2007 : cadre de sondage et échantillonnage. Numéro de commande : 338-0052 / Office Fédéral de la Statistique. Neuchâtel, 2008. – BFS Methodenbericht. – URL <https://www.bfs.admin.ch/bfs/fr/home/services/recherche/rapports-methodologiques.assetdetail.344506.html>
- [OFS: Renaud A., Panchard Ch., Potterat J. 2008] OFS: RENAUD A., PANCHARD CH., POTTERAT J.: Statistique de l'emploi. Révision 2007 : méthodes d'estimation. Numéro de commande : 338-0055 / Office Fédéral de la Statistique. Neuchâtel, 2008. – BFS Methodenbericht. – URL <https://www.bfs.admin.ch/bfs/fr/home/services/recherche/rapports-methodologiques.assetdetail.344658.html>
- [Schouten B. 2017] SCHOUTEN B., Wagner J.: *Adaptive Survey Design*. New York : Chapman & Hall/CRC, 2017

Methodenberichte der Sektion Statistische Methoden des BFS
Rapports de méthodes de la section méthodes statistiques de l'OFS
Methodology reports published by the FSO's Statistical Methods Section

- Potterat, J., Assoulin, D., Nicoletti, J.-M. (2019). Beschäftigungsstatistik BESTA. Revision 2015: Stichprobenrahmen, Stichprobenplan und Hochrechnung. Bestellnummer: 338-0079
- Massiani, A. (2017). Estimation de la couverture du nouveau recensement en Suisse en 2012. Numéro de commande : 338-0078
- Nedyalkova, D., Assoulin, D. (2017). Construction of full-time equivalents for the Swiss structural business statistics. Order number: 338-0077
- Ferster, M. (2016). Energieverbrauchsstatistik EVS2014 - Stichprobe, Hochrechnung und Vergleichbarkeit mit der EVS2013. Bestellnummer: 338-0076
- Panchard, C. (2014). Enquête sur la participation aux activités culturelles 2008. Plan de sondage et estimations. Bestellnummer: 338-0073
- Assoulin, D. (2014). Wertschöpfungsstatistik. Revision 2009: Statistische Datenaufbereitung und Hochrechnung. Bestellnummer: 338-0072
- Wilhelm, M. (2014). Echantillonnage boule de neige. La méthode de sondage déterminé par les répondants. Numéro de commande: 338-0071
- Assoulin, D. (2013). Wertschöpfungsstatistik. Revision 2009: Stichprobenrahmen und Stichprobenplan. Bestellnummer: 338-0070
- Ferster, M. (2013). EVS I - Energieverbrauchsstatistik 2002 bis 2007: Stichprobenplan und Hochrechnung. Bestellnummer: 338-0069
- Assoulin, D. (2013). Zusatzerhebung für die landwirtschaftliche Betriebszählung 2010: Stichprobenplan und Hochrechnung. Bestellnummer: 338-0068
- Potterat, J. (2012). Use of conversion keys for NOGA2002-2008. Order number: 338-0067-05
- Andrade, B., Salamin P.-A. (2012). Enquête sur la situation sociale et économique des étudiant-e-s des hautes écoles suisses 2009. Cadre de sondage, plan d'échantillonnage et méthodes d'estimation. Numéro de commande: 338-0066
- Potterat, J., Panchard, C., Kilchmann, D. (2012). Umweltschutzausgaben der Unternehmen 2009 (UWSA2009). Stichprobenplan, Einsetzungen, Gewichtung und Schätzverfahren. Bestellnummer: 338-0065
- Potterat, J. (2012). Benutzung der Umsteigeschlüssel NOGA 2002-2008. Bestellnummer: 338-0064
- Potterat, J. (2011). Kosten und Nutzen der Berufsbildung aus Sicht der Betriebe im Jahr 2009 (KNBB09). Stichprobenplan, Gewichtung und Schätzverfahren. Bestellnummer: 338-0063
- Kilchmann, D., Potterat, J., Genoud, S. (2011). Gütertransporterhebung 2008. Stichprobenplan, Datenaufbereitung, Gewichtung und Schätzverfahren. Bestellnummer: 338-0062
- Graf, E. (2010). Enquête suisse sur la santé 2007. Plan d'échantillonnage, pondérations et analyses pondérées des données. Numéro de commande: 338-0061
- Eichenberger P., Hulliger B., Potterat J. (2010). Describing the Anticipated Accuracy of the Swiss Population Survey. Order number: 338-0060
- Graf, E. (2010). Étude empirique de l'attrition du Panel Suisse de Ménages : vers une caractérisation du profil du non-répondant. Numéro de commande: 338-0059
- Graf, E. (2009). Weightings of the Swiss Household Panel: SHP_I wave 9, SHP_II wave 4, SHP_I et SHP_II combined. Order number: 338-0058
- Graf, E. (2009). Pondérations du Panel Suisse de Ménages: PSM_I vague 9, PSM_II vague 4, PSM_I et PSM_II combinés. Numéro de commande: 338-0057-05

Qualité, L., Tillé, Y. (2009). Estimation de la précision d'évolutions dans l'enquête sur la valeur ajoutée. Numéro de commande: 338-0056

Renaud, A., Panchard, C. et Potterat, J. (2008). Statistique de l'emploi. Révision 2007 : méthodes d'estimation. Numéro de commande: 338-0055

Graf, E. (2008). Pondérations du PSM. PSM_I vague 8, PSM_II vague 3, PSM_I et PSM_II combinés. Numéro de commande: 338-0054

Andrade, B., Graf, M. (2008). Enquête suisse sur la structure des salaires 2006. Aspects méthodologiques du modèle des salaires SSalarium". Numéro de commande: 338-0053

Renaud, A. (2008). Statistique de l'emploi. Révision 2007 : cadre de sondage et échantillonnage. Numéro de commande: 338-0052

Graf, E. (2008). Pondérations du SILC pilote. SILC_I vague 2, SILC_II vague 1, SILC_I et SILC_II combinés. Numéro de commande: 338-0051

Kilchmann, D. (2008). Statistik der sozialmedizinischen Institutionen 1999-2004 und Krankenhausstatistik 1999-2002. Einsetzungen für fehlende Daten. Bestellnummer: 338-0050

Renaud, A. (2008). Technologies de l'information et de la communication. Estimations sur la base de la statistique de la valeur ajoutée. Numéro de commande: 338-0049

Assoulin, D. (2007). Wertschöpfungsstatistik. Einsetzungsversuche für fehlende Antworten grosser Unternehmen. Bestellnummer: 338-0048

Kilchmann, D. (2007). Beherbergungsstatistik Campingplätze. Stichprobenrahmen und Schätzverfahren 2005/06. Bestellnummer: 338-0047

Gabler, S., Häder, S. (2007). Haushalts- und Personenerhebungen. Machbarkeit von Random Digit Dialing in der Schweiz. Bestellnummer: 338-0046

Ferrez, J., Graf, M. (2007). Enquête suisse sur la structure des salaires. Programmes R pour l'intervalle de confiance de la médiane. Numéro de commande: 338-0045

Renaud, A. (2007). Harmonisation de la scolarité obligatoire en Suisse (HarmoS). Design général de l'enquête et échantillon des écoles. Numéro de commande: 338-0044

Potterat, J. (2007). Betriebszählung 2005. Statistische Methoden zur Schätzung der provisorischen Ergebnisse. Bestellnummer: 338-0043

Hulliger, B. (2006). Umweltschutzausgaben der Unternehmen 2003, Stichprobenplan, Datenaufbereitung und Schätzverfahren. Bestellnummer: 338-0042

Renfer, J.-P. (2006). Enquête sur les chiffres d'affaires du commerce de détail. Plan d'échantillonnage et méthodes d'estimation. Numéro de commande: 338-0041

Salamin, P.-A. (2006). Statistique de l'aide sociale dans le domaine de l'asile. Plan de sondage et extrapolations pour l'enquête pilote 2005. Numéro de commande: 338-0040

Renaud, A. (2006). Statistique suisse des bénéficiaires de l'aide sociale. Pondération des communes 2004. Numéro de commande: 338-0039

Graf, M. (2006). Swiss Earnings Structure Survey 2002-2004. Compositional data in a stratified two-stage sample: Analysis and precision assessment of wage components. Order number: 338-0038

Potterat, J. (2006). Pensionskassenstatistik 2004. Statistische Methoden zur Schätzung der provisorischen Ergebnisse. Bestellnummer: 338-0037

Potterat, J. (2006). Kosten und Nutzen der Berufsbildung aus Sicht der Betriebe im Jahr 2004. Stichprobenplan, Gewichtung und Schätzverfahren. Bestellnummer: 338-0036

Kilchmann, D. (2006). Vierteljährliche Wohnbaustatistik. Stichprobenplan, statistische Datenaufbereitung und Schätzverfahren 2005. Bestellnummer: 338-0035

Kilchmann, D. (2006). Erhebung über Forschung und Entwicklung in der schweizerischen Privatwirtschaft 2004. Bereinigung der Stichprobe, Ersatz fehlender Werte und Schätzverfahren. Bestellnummer: 338-0034

- Kilchmann, D., Eichenberger, P., Potterat, J. (2005). Volkszählung 2000. Statistische Einsetzungsverfahren Band 2. Bestellnummer: 338-0033
- Kilchmann, D., Eichenberger, P., Potterat, J. (2005). Volkszählung 2000. Statistische Einsetzungsverfahren Band 1. Bestellnummer: 338-0032
- Graf, M., Matei, A. (2005). Enquête suisse sur la structure des salaires 2002. La précision du salaire brut standardisé médian. Numéro de commande: 338-0031
- Graf, E., Renfer, J.-P. (2005). Enquête suisse sur la santé 2002. Plan d'échantillonnage, pondération et estimation de la précision. Numéro de commande: 338-0030
- Potterat, J. (2005). Mietpreis-Strukturerhebung 2003. Gewichtung und Schätzverfahren. Bestellnummer: 338-0029
- Potterat, J. (2005). Landwirtschaftliche Betriebszählung 2003. Schätzverfahren für die Zusatzerhebung. Bestellnummer: 338-0028
- Renaud, A. (2004). Coverage estimation for the Swiss population census 2000. Estimation methodology and results. Order number: 338-0027
- Kilchmann, D. (2004). Revision des Schweizerischen Lohnindex. Schätzmethoden der Lohnindices und deren Varianzschätzer. Bestellnummer: 338-0026
- Graf, M. (2004). Enquête suisse sur la structure des salaires 2002. Plan d'échantillonnage et extrapolation pour le secteur privé. Numéro de commande: 338-0025
- Renaud, A. (2004). Analyse de données d'enquêtes. Quelques méthodes et illustration avec des données de l'OFS. Numéro de commande 338-0024
- Renaud, A., Potterat, J. (2004). Estimation de la couverture du recensement de la population de l'an 2000. Echantillon pour l'estimation de la sous-couverture (P-sample) et qualité du cadre de sondage des bâtiments. Numéro de commande: 338-0023
- Graf, M. (2004). Fusion de données. Etude de faisabilité. Numéro de commande: 338-0022
- Potterat, J. (2003). Mietpreis-Strukturerhebung 2003. Entwicklung des Stichprobenplans und Ziehung der Stichprobe. Bestellnummer: 338-0021
- Potterat, J. (2003). Landwirtschaftliche Betriebszählung 2003. Stichprobenplan der Zusatzerhebung. Bestellnummer: 338-0020.
- Renaud, A. (2003). Estimation de la couverture du recensement de la population de l'an 2000. Echantillon pour l'estimation de la sur-couverture (E-sample). Numéro de commande: 338-0019
- Hulliger, B. (2003). Bereinigung der Stichprobe, Ersatz fehlender Werte und Schätzverfahren. Erhebung über F+E in der schweizerischen Privatwirtschaft 2000. Bestellnummer: 338-0018
- Renfer, J.-P. (2003). Enquête 2000 sur la recherche et le développement dans l'économie privée en Suisse. Plan d'échantillonnage. Numéro de commande: 338-0017
- Potterat, J. (2003). Kosten und Nutzen der Berufsbildung aus Sicht der Betriebe. Schätzverfahren. Bestellnummer: 338-0016
- Graf, M., Matei, A. (2003). Stratégie de choix des modèles de désaisonnalisation. Application aux séries de l'emploi total. Numéro de commande: 338-0015
- Potterat, J., Salamin, P.A. (2002). Betriebszählung 2001. Methoden für die Datenbereinigung. Bestellnummer: 338-0014
- Renaud, A. (2002). Programme international pour le suivi des acquis des élèves (PISA). Plans d'échantillonnage pour PISA 2000 en Suisse. Numéro de commande: 338-0013
- Renfer, J.-P. (2002). Enquête 2001 sur les coûts et l'utilité de la formation des apprentis du point de vue des établissements. Plan d'échantillonnage. Numéro de commande: 338-0012
- Potterat, J., Salamin, P.A. (2002). Betriebszählung 2001. Stichprobenplan und Schätzverfahren für die provisorischen Ergebnisse. Bestellnummer: 338-0011

- Graf, M. (2002). Enquête suisse sur la structure des salaires 2000. Plan d'échantillonnage, pondération et méthode d'estimation pour le secteur privé. Numéro de commande: 338-0010
- Renaud, A., Eichenberger P. (2002). Estimation de la couverture du recensement de la population de l'an 2000. Procédure d'enquête et plan d'échantillonnage de l'enquête de couverture. Numéro de commande: 338-0009
- Kilchmann, D., Hulliger, B. (2002). Stichprobenplan für die Obstbaumzählung 2001. Bestellnummer: 338-0008
- Graf, M. (2002). Passage du concept établissement au concept entreprise. Numéro de commande: 338-0007
- Salamin, P.A. (2001). La technique de la double enquête pour la statistique du transport routier de marchandise. Numéro de commande: 338-0006
- Peters, R., Renfer, J.-P. et Hulliger, B. (2001). Statistique de la valeur ajoutée 1997-1998. Procédure d'extrapolation des données. Numéro de commande: 338-0005
- Potterat, J., Hulliger, B. (2001). Schätzung der Sägereiproduktion mit der Sägerei-Erhebung PAUL. Bestellnummer: 338-0004
- Graf, M. (2001). Désaisonnalisation. Aspects méthodologiques et application à la statistique de l'emploi. Numéro de commande: 338-0003
- Hüsler, J., Müller, S. (2001). Schlussbericht Betriebszählung 1995 (BZ 95), Mehrfach imputierte Umsatzzahlen. Bestellnummer: 338-0002
- Renaud, A. (2001). Statistique suisse des bénéficiaires de l'aide sociale. Plan d'échantillonnage des communes. Numéro de commande: 338-0001
- Hulliger, B., Eichenberger, P. (2000). Stichprobenregister für Haushalterhebungen: Umstellung auf Telefonnummern ohne Namen und Adressen, Abläufe für Erstellung und Stichprobenziehung. Bestellnummer: 338-0000

