

2023

06

Industrie und
Dienstleistungen

Neuchâtel 2023

Schätzung des Umsatzes in der Statistik der Unternehmensgruppen (STAGRE)

Methodische Grundlagen



Schweizerische Eidgenossenschaft
Confédération suisse
Confederazione Svizzera
Confederaziun svizra

Eidgenössisches Departement des Innern EDI
Bundesamt für Statistik BFS

Themenbereich «Industrie und Dienstleistungen»

Aktuelle themenverwandte Publikationen

Fast alle vom BFS publizierten Dokumente werden auf dem Portal www.statistik.ch gratis in elektronischer Form zur Verfügung gestellt. Gedruckte Publikationen können bestellt werden unter der Telefonnummer +41 58 463 60 60 oder per E-Mail an order@bfs.admin.ch.

Unternehmensdemografie (UDEMÖ): Analysen der Daten von 2013 bis 2020, Neuchâtel 2022, 12 Seiten, Bestellnummer: 1783-2000

Statistik der Unternehmensgruppen (STAGRE): Portrait der Unternehmensgruppen in der Schweiz 2014–2021, Neuchâtel 2022, 12 Seiten, Bestellnummer: 1844-2100

Porträt der Schweizer KMU, 2011–2020, Neuchâtel 2022, 12 Seiten, Bestellnummer: 1661-2000

Buchhaltungsergebnisse schweizerischer Unternehmen: Geschäftsjahre 2019–2020, Neuchâtel 2022, 104 Seiten, Bestellnummer: 029-2001

Arbeitsmarktindikatoren 2021, Neuchâtel 2021, 84 Seiten, Bestellnummer: 206-2101

Themenbereich «Industrie und Dienstleistungen» im Internet

www.statistik.ch → Statistiken finden → 06 – Industrie, Dienstleistungen

Schätzung des Umsatzes in der Statistik der Unternehmensgruppen (STAGRE)

Methodische Grundlagen

Redaktion Marius Ley, BFS
Herausgeber Bundesamt für Statistik (BFS)

Neuchâtel 2023

Herausgeber: Bundesamt für Statistik (BFS)

Auskunft: Marius Ley, BFS, Tel. +41 58 463 66 13,
marius.ley@bfs.admin.ch

Redaktion: Marius Ley, BFS

Reihe: Statistik der Schweiz

Themenbereich: 06 Industrie und Dienstleistungen

Originaltext: Deutsch

Layout: Publishing und Diffusion PUB, BFS
Sektion WSA, BFS

Grafiken: Sektion WSA, BFS

Online: www.statistik.ch

Print: www.statistik.ch
Bundesamt für Statistik, CH-2010 Neuchâtel,
order@bfs.admin.ch, Tel. +41 58 463 60 60
Druck in der Schweiz

Copyright: BFS, Neuchâtel 2023
Wiedergabe unter Angabe der Quelle
für nichtkommerzielle Nutzung gestattet

BFS-Nummer: 2239-2300

ISBN: 978-3-303-06342-2

Inhaltsverzeichnis

1	Kontext und Datenlage	5
1.1	Statistik der Unternehmensgruppen (STAGRE) und die Variable «Umsatz»	5
1.2	Datenlage	6
2	Anforderungen und Spezifikation des verwendeten Modells	11
2.1	Ausgangslage zur Modellierung	11
2.2	Verwendetes Modell	11
3	Evaluation der verschiedenen Modellvarianten	13
3.1	Grundlegende Herausforderungen	13
3.2	Kriterien und Grundsätze	14
3.3	Erste Etappe (Basisvarianten)	16
	Grundsätzliche Überlegungen	16
	Vorgehen und Resultate	18
3.4	Zweite Etappe (Detailvarianten)	21
	Grundsätzliche Überlegungen	21
	Vorgehen und Resultate	25
	Grafische Darstellung und punktuelle Korrekturen	26
4	Abschliessende Bemerkungen und Ausblick	31
4.1	Abschliessende Bemerkungen	31
4.2	Ausblick	31
5	Anhang	33
5.1	Detailmodelle: Ausführliche Erläuterungen	33
6	Technischer Anhang	37

1 Kontext und Datenlage

1.1 Statistik der Unternehmensgruppen (STAGRE) und die Variable «Umsatz»

Im November 2018 veröffentlichte das Bundesamt für Statistik (BFS) erstmals die Statistik der Unternehmensgruppen (STAGRE) und ergänzte damit das Angebot der strukturellen Unternehmensstatistik. Zu den von der STAGRE erfassten Einheiten¹ gehören einerseits Mitglieder von Unternehmensgruppen mit einem Gruppenoberhaupt im Ausland, so dass die Aktivitäten und Bedeutung von ausländisch kontrollierten Unternehmen in der Schweiz beziffert werden können. Andererseits erfasst diese Statistik auch die Unternehmensgruppen unter inländischer Kontrolle, um eine Einschätzung des Stellenwerts der multinationalen Unternehmen und der Unternehmensgruppen im weiteren Sinne im Schweizer Wirtschaftsgefüge zu ermöglichen. Seit ihrem ersten Erscheinen beinhaltet die STAGRE Informationen zur Anzahl der Einheiten und der Beschäftigten, seit 2019 auch zum internationalen Warenhandel. Die Resultate der STAGRE werden jährlich aktualisiert und reichen zurück bis zum Referenzjahr 2014.

Um die Analysemöglichkeiten zu erweitern, und um den internationalen Erfordernissen im Gebiet der «Foreign Affiliates Statistics» (FATS)² zu entsprechen, wurde entschieden, für die Unternehmen der STAGRE für die gesamte Zeitreihe auch Resultate für den Umsatz bereitzustellen, die aggregiert nach verschiedenen Dimensionen (wie Grössenklasse, Art der Gruppe, Branche, Sitzland etc.) verfügbar sein sollten. Dies ist auch im Sinne der thematischen Schwerpunkte des aktuellen Statistischen Mehrjahresprogrammes.³ Die Bereitstellung dieser Kenngrösse war allerdings nicht trivial, da für den Umsatz – im Gegensatz zu den weiter oben erwähnten Variablen – keine statistische Quelle existiert, aus der sich mittels Datenverknüpfungen unmittelbar verwertbare Resultate ableiten lassen. Die Analyse der potentiell in Frage kommenden Quellen zeigte hingegen, dass die Verknüpfung mehrerer Quellen und, darauf aufbauend, statistische Modellierungen einen vielversprechenden Ansatz darstellen, um zuverlässige, nach den oben erwähnten Dimensionen aufschlüsselbare Resultate zu produzieren. Im November 2020 wurden solche

modellbasierten Resultate für den Umsatz erstmals publiziert. Die Datenreihe wurde seither jährlich aktualisiert und das dafür verwendete Verfahren stetig verfeinert. Sowohl das statistische Modell, das diesen Resultaten zugrunde liegt, als auch die Vorgehensweise zur Evaluation dieses Modells in Bezug auf mögliche Alternativvarianten wurden seit der ersten Veröffentlichung jedoch keinen radikalen Veränderungen unterworfen.

Im Sinne einer grösstmöglichen Transparenz bezüglich der Produktion der Variable «Umsatz» erläutert der vorliegende Bericht die Ausgangslage in Bezug auf die verfügbaren Daten (Abschnitt 1.2), das schlussendlich verwendete Modell (Kapitel 2) sowie die Kriterien und Entscheidungen, die definiert und getroffen werden mussten, um diese Variable in der geforderten Qualität in der STAGRE zur Verfügung zu stellen (Kapitel 3). In erster Linie richtet sich dieser Bericht an ein technisch versiertes, mit den Begriffen der statistischen Modellierung vertrautes Publikum, da eine Vielzahl von technischen Details behandelt werden. Beschrieben wird der Stand der Modellierung, wie sie für die Schätzung der jüngsten, im November 2022 veröffentlichten Resultate der STAGRE verwendet wurde. Für die Erarbeitung sowohl des Modells als auch des vorliegenden Berichts konnte auf eine enge Begleitung durch die Sektion Statistische Methoden (METH) des Bundesamtes für Statistik zurückgegriffen werden.

¹ Der Begriff «Einheit» wird im vorliegenden Text grundsätzlich synonym zum Begriff «Unternehmen» verwendet.

² Für Statistiken zu multinational tätigen Unternehmen sind die auf europäischer Ebene relevanten Standards, an denen sich auch die STAGRE orientiert, im Rahmen der FATS definiert. Siehe [https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Glossary:Foreign_affiliates_statistics_\(FATS\)](https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Glossary:Foreign_affiliates_statistics_(FATS))

³ Statistisches Mehrjahresprogramm des Bundes 2020–2023, Neuchâtel 2020. Siehe Abschnitt 3.9, Wirtschaft und öffentliche Finanzen, Schwerpunkte: «Die makroökonomischen Daten werden dahingehend analysiert, dass sie die Entwicklungen im Zusammenhang mit der Globalisierung (...) besser aufzeigen können. Auf mikroökonomischer Ebene stehen die multinationalen Unternehmen im Zentrum (...).»

1.2 Datenlage

Potentielle Quellen für den Umsatz

Bei der Erarbeitung der STAGRE war klar, dass der Grossteil der erforderlichen strukturellen Variablen – wie Anzahl Unternehmen, Arbeitsstätten, Beschäftigte sowie die Branchenzugehörigkeit (NOGA-Code) – der Statistik der Unternehmensstruktur (STATENT) entnommen werden kann.⁴ Dies gilt jedoch nicht für die Variable «Umsatz», da diese in der STATENT nicht enthalten ist. Mit der **Wertschöpfungsstatistik (WS)** des BFS steht zwar eine potentielle Datenquelle für diese Variable zur Verfügung. Bei der WS handelt es sich jedoch um eine Stichprobenerhebung mit dem Zweck, betriebswirtschaftliche Kennzahlen auf Ebene der NOGA-Abteilungen⁵ für die Grundgesamtheit aller in der Schweiz ansässigen Unternehmen mit mindestens 3 Beschäftigten bereitzustellen. Eine Berechnung von Kennzahlen der WS für eine Untergruppe dieser Grundgesamtheit – im konkreten Fall für all jene Einheiten, die Teil einer Unternehmensgruppe sind – ist nicht ohne weiteres möglich, da der Stichprobenplan nicht für eine solche Verwendung konzipiert wurde.

Als Alternative zur WS wäre deshalb theoretisch denkbar, die durch die Eidgenössische Steuerverwaltung (ESTV) zum Zweck der Entrichtung der **Mehrwertsteuer (MWST)** erhobenen Daten zu verwenden. Die darin enthaltenen Umsätze erfassen sämtliche steuerpflichtigen Unternehmen in der Schweiz und stehen im BFS als Grundlage für verschiedene Schätzungen von Statistiken zur Verfügung. Eine Analyse dieser Daten ergab jedoch, dass sie aus mehreren Gründen nicht für eine unmittelbare statistische Verwendung in Frage kommen:

- Die Definition des für die MWST relevanten Umsatzes entspricht im Gegensatz zur WS nicht exakt jener der strukturellen Unternehmensstatistik (siehe Kasten); so sind etwa in den MWST-Umsätzen die «sonstigen betrieblichen Erträge» enthalten, die gemäss statistischer Definition auszuschliessen sind.
- Gewisse Umsatzkomponenten, insbesondere Exporte, werden von manchen Unternehmen in der Deklaration weggelassen (da diese Komponenten nicht relevant sind für die Bemessung der Steuer).

- Für Unternehmensgruppen besteht die Option, zwecks administrativer Erleichterung die MWST gesamthaft als Gruppe statt separat für jede einzelne Einheit abzurechnen. Für diese Einheiten stehen für statistische Zwecke keine unmittelbar (oder «original») erhobene Umsatzdaten zur Verfügung. Zwar wurde im BFS 2015 ein Schätzverfahren entwickelt (und 2022 verfeinert), das die mittels Gruppenbesteuerung erhobenen Umsätze auf die einzelnen involvierten Einheiten verteilt; die Zuverlässigkeit dieser «verteilten» Umsätze ist jedoch naturgemäss geringer, als dies bei direkt erhobenen Umsätzen der Fall wäre.
- Um das Steuergeheimnis zu wahren, ist gemäss einer Vereinbarung zwischen der ESTV und dem BFS eine direkte Verwendung der MWST-Umsätze als statistische Variable nicht zulässig. Einer Verwendung dieser Variablen für Schätzungen von Statistiken steht hingegen nichts im Wege.

Umsatz: Definition von Eurostat im Rahmen der Strukturellen Unternehmensstatistik:

Der **Umsatz**, im Kontext der strukturellen Unternehmensstatistik, umfasst die von der Erhebungseinheit während des Berichtszeitraums insgesamt in Rechnung gestellten Beträge, die den Verkäufen von Waren und Dienstleistungen an Dritte entsprechen.

Der Umsatz schliesst ein:

- alle Steuern und Abgaben, die auf den von der Einheit in Rechnung gestellten Waren oder Dienstleistungen liegen, mit Ausnahme der Mehrwertsteuer (MWST) und ähnlicher abziehbarer Abgaben, die in direktem Zusammenhang zum Umsatz stehen;
- alle berechneten Nebenkosten (Transport, Verpackung usw.), die an die Kunden weitergegeben werden, selbst wenn diese Kosten getrennt in Rechnung gestellt werden.

Preisnachlässe, Rabatte und Skonti sowie der Wert der zurückgegebenen Verpackung sind abzuziehen.

Ausgenommen sind:

- Erträge, die in den Unternehmensabschlüssen als sonstige betriebliche Erträge, Finanzerträge oder ausserordentliche Erträge ausgewiesen sind;
- vom Staat oder der Europäischen Union erhaltene Betriebssubventionen.

[https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Glossary:Turnover_-_SBS_\(in_English\)](https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Glossary:Turnover_-_SBS_(in_English))

⁴ In der STATENT sind alle in der Schweiz ansässigen Unternehmen enthalten, die im jeweiligen Referenzjahr im Monat Dezember mindestens eine Person beschäftigten, für die AHV-Beiträge entrichtet wurden. Da die STAGRE auch Einheiten umfasst, die über keine Beschäftigten verfügen, ist sie nicht eine strikte Untergruppe der STATENT; jedoch werden in der STAGRE für alle Einheiten mit Beschäftigten die Informationen grundsätzlich der STATENT entnommen.

⁵ Die Allgemeine Systematik der Wirtschaftszweige (NOGA) ist die für die öffentliche Statistik der Schweiz gebräuchliche Branchenklassifizierung. Eine Definition der NOGA-Abteilungen findet sich in Tabelle TA4 im Anhang.

Zusammenfassend lässt sich somit sagen: Für die Variable «Umsatz» in der STAGRE steht mit der WS eine Quelle zur Verfügung, die bezüglich Definition und Qualität den Anforderungen entspricht, deren Abdeckung jedoch auf Grund des Stichprobencharakters unzureichend ist. Umgekehrt haben die Umsätze der MWST eine ziemlich gute Abdeckung, entsprechen aber nicht der gewünschten Definition und Qualität. Eine Möglichkeit, um diese Einschränkungen zu überwinden und die spezifischen Vorteile jeweils beider Quellen zu nutzen, besteht in einer Verknüpfung der beiden Quellen auf Ebene der Unternehmen und einer **statistischen Modellierung**: Es wird also ein Modell spezifiziert, das für jedes Unternehmen einen Schätzwert (synonym: eine Vorhersage) für die Zielvariable «Umsatz» liefert, der möglichst nahe am Wert gemäss der qualitativ bevorzugten Quelle (also der WS) liegt. Eine solche Vorhersage kann für all jene Einheiten berechnet werden, für welche die in das Modell einflussenden Hilfsvariablen⁶ verfügbar sind. Somit bieten sich die MWST-Umsätze als potentielle Hilfsvariable an. Für Modellvorhersagen von ausreichender Qualität sind die folgenden **Bedingungen** erforderlich:

- a) Die Hilfsvariablen müssen in ausreichendem Mass mit der Zielvariablen korrelieren.
- b) Es muss eine genügend grosse Menge an Beobachtungen mit vorhandenen Werten sowohl für die Ziel- als auch für die Hilfsvariablen zur Verfügung stehen, auf Grund derer sich das Modell schätzen lässt.
- c) Die Einheiten, für welche die Vorhersagen verwendet werden sollen, dürfen sich in struktureller Hinsicht nicht wesentlich von den Einheiten unterscheiden, mittels derer das Modell geschätzt wurde.

Eine einfache Korrelationsanalyse deutet darauf hin, dass Bedingung (a) erfüllt ist.⁷ Für eine Beurteilung der Bedingungen (b) und (c) und – damit verbunden – eine Einschätzung darüber, ob für alle Einheiten der STAGRE Modellvorhersagen berechnet werden können oder nur für eine Teilmenge dieser Einheiten, ist es erforderlich, die Datenlage in Bezug auf das Vorhandensein von Werten für die Ziel- und Hilfsvariablen genauer zu analysieren. Im Zuge dieser Analysen soll auch untersucht werden, inwiefern als weitere Hilfsvariablen Beschäftigtenzahl, Vollzeitäquivalente Beschäftigung (VZÄ) und/oder Lohnsumme gemäss STATENT/AHV geeignet sein könnten.^{8,9}

⁶ Grundsätzlich ist die Verwendung sowohl von einer einzigen als auch von mehreren Hilfsvariablen denkbar. Da die schlussendlich verwendete Variante mehr als eine Hilfsvariable beinhaltet, wird im Text durchwegs der Plural verwendet.

⁷ Der Korrelationskoeffizient zwischen MWST- und WS-Umsatz beträgt 0.83, zwischen den logarithmierten Werten der beiden Variablen sogar 0.97 (Berechnungsgrundlage: Einheiten mit verfügbaren Daten, Referenzjahr 2020).

⁸ Der Korrelationskoeffizient zwischen der Beschäftigtenzahl gemäss STATENT und dem Umsatz gemäss WS beträgt nur 0.19, zwischen den logarithmierten Werten der beiden Variablen jedoch 0.80 (wiederum für alle Einheiten mit verfügbaren Daten, Referenzjahr 2020). Für VZÄ gemäss STATENT und für die AHV-Lohnsumme resultieren leicht höhere Korrelationswerte mit dem WS-Umsatz.

Der Einbezug dieser oder einer dieser Hilfsvariablen könnte in einer geeigneten funktionalen Spezifikation somit sinnvoll sein.

⁹ Hauptquelle für diese Kandidaten von Hilfsvariablen sind die AHV-Ausgleichskassen. In der Folge gibt es eine sehr grosse Übereinstimmung zwischen den Populationen derjenigen Einheiten, für welche Werte dieser Variablen jeweils verfügbar sind, wobei die Abdeckung der Lohnsumme leicht bessere ist als jene der Beschäftigtenzahl und der VZÄ. Der Einfachheit halber beschränken sich die hier vorgenommenen Analysen auf die potentielle Hilfsvariable «Anzahl Beschäftigte»; die Erkenntnisse daraus lassen sich problemlos auf die AHV-Lohnsumme und die STATENT-VZÄ übertragen.

Tabelle T1a: Verteilung der Einheiten der STAGRE-Grundgesamtheit, Referenzjahr 2020

Einheiten in Unternehmensgruppen: Anzahl der Einheiten ...	Mit Daten verfügbar aus									
	STATENT und MWST		nur STATENT		nur MWST		(keine dieser beiden)		Total	
... in der Nettostichprobe der WS	5'639	9.8%	65	0.1%	8	0.0%	1	0.0%	5'713	10.0%
... ausserhalb Nettostichprobe, aber im Rahmen der WS	13'364	23.3%	621	1.1%	0	0.0%	0	0.0%	13'985	24.4%
... in WS-relevanter NOGA/RF, aber mit < 3 Beschäftigten	4'838	8.4%	1'255	2.2%	11'052	19.3%	7'131	12.4%	24'276	42.4%
... ausserhalb WS-relevanter NOGA/Rechtsform	3'144	5.5%	1'415	2.5%	2'192	3.8%	6'581	11.5%	13'332	23.3%
Total	26'985	47.1%	3'356	5.9%	13'252	23.1%	13'713	23.9%	57'306	100.0%

Tabelle T1b: Verteilung der Beschäftigten der STAGRE-Grundgesamtheit, Referenzjahr 2020

Einheiten in Unternehmensgruppen: Anzahl der Beschäftigten in Einheiten...	Mit Daten verfügbar aus									
	STATENT und MWST		nur STATENT		nur MWST		(keine dieser beiden)		Total	
... in der Nettostichprobe der WS	1'284'869	67.1%	5'612	0.3%	0	0.0%	0	0.0%	1'290'481	67.4%
... ausserhalb Nettostichprobe, aber im Rahmen der WS	294'442	15.4%	12'328	0.6%	0	0.0%	0	0.0%	306'770	16.0%
... in WS-relevanter NOGA/RF, aber mit < 3 Beschäftigten	6'795	0.4%	1'584	0.1%	0	0.0%	0	0.0%	8'379	0.4%
... ausserhalb WS-relevanter NOGA/Rechtsform	296'875	15.5%	12'236	0.6%	0	0.0%	0	0.0%	309'111	16.1%
Total	1'882'981	98.3%	31'760	1.7%	0	0.0%	0	0.0%	1'914'741	100.0%

Tabelle T1c: Verteilung des MWST-Umsatzes der STAGRE-Grundgesamtheit, Referenzjahr 2020

Einheiten in Unternehmensgruppen: Umsatz gemäss MWST (in Mio. Franken) in Einheiten...	Mit Daten verfügbar aus									
	STATENT und MWST		nur STATENT		nur MWST		(keine dieser beiden)		Total	
... in der Nettostichprobe der WS	2'012'999	67.2%	0	0.0%	184	0.0%	0	0.0%	2'013'183	67.2%
... ausserhalb Nettostichprobe, aber im Rahmen der WS	447'671	15.0%	0	0.0%	0	0.0%	0	0.0%	447'671	15.0%
... in WS-relevanter NOGA/RF, aber mit < 3 Beschäftigten	35'697	1.2%	0	0.0%	92'616	3.1%	0	0.0%	128'314	4.3%
... ausserhalb WS-relevanter NOGA/Rechtsform	395'677	13.2%	0	0.0%	9'532	0.3%	0	0.0%	405'209	13.5%
Total	2'892'044	96.6%	0	0.0%	102'332	3.4%	0	0.0%	2'994'376	100.0%

Vorhandensein der Variablen

Die Analyse zum Vorhandensein von Werten für die Ziel- und Hilfsvariablen bezieht sich der Einfachheit halber ausschliesslich auf die Resultate des Referenzjahres 2020, dargestellt in den Tabellen T1 und T2.¹⁰ Diese enthalten, für die Grundgesamtheit der STAGRE (alle Unternehmen, die Teil einer Unternehmensgruppe sind), die Anzahl der Einheiten sowie deren Beschäftigtenzahl und Umsatz, aufgeschlüsselt einerseits nach dem Status in Bezug auf die Wertschöpfungsstatistik (WS) und andererseits nach dem Vorhandensein von Werten für die beiden potentiellen Hilfsvariablen STATENT-Beschäftigte und MWST-Umsatz. Daraus ergeben sich die folgenden Erkenntnisse:

- *Jeweils oberste Zeile:* Nur 10.0% der Einheiten der STAGRE (5 713 an der Zahl) sind in der **Nettostichprobe der WS** und verfügen somit über erhobene WS-Umsätze (Tabelle T1a).
- *Jeweils unterste Zeile:* In der STAGRE sind Einheiten mit **NOGA/Rechtsform**-Kombinationen enthalten, die **für den Rahmen der WS als nicht relevant** definiert wurden (23.3% an der Zahl; Tabelle T1a). Eine zuverlässige Schätzung des WS-Umsatzes für diese Subpopulation ist nicht möglich (Verletzung der oben unter (c) genannten Bedingung). Die Bedeutung dieser Einheiten ist allerdings geringer, wenn deren Anteile gewichtet nach Beschäftigung (16.1%; Tabelle T1b) und nach MWST-Umsatz (13.5%; Tabelle T1c) herangezogen werden. Zudem zeigte eine Analyse nach Branchen, dass hier Einheiten der NOGA-Abteilungen 64–66 (Finanz- und Versicherungsdienstleistungen) und 86 (Gesundheitswesen) zusammengenommen einen nach Beschäftigung oder MWST-Umsatz gewichteten Anteil von mehr als 96% ausmachen – also Abteilungen, die praktisch vollumfänglich aus dem Rahmen der WS ausgeschlossen sind.¹¹ Fazit: Ein Verzicht auf Modellvorhersagen für diese Einheiten ist unausweichlich, schmälert den Nutzen des Modellierungsansatzes aber nicht übermässig.¹²

- *Jeweils zweite und dritte Zeile:* Etwa zwei Drittel (66.8%) der Einheiten der STAGRE haben eine für die WS relevante NOGA/Rechtsform-Kombination, verfügen aber über keine erhobenen WS-Umsätze. Diese zwei Drittel setzen sich zusammen aus

- *zweite Zeile:* 24.4% Einheiten mit drei oder mehr Beschäftigten und somit zum «ex post» konstruierten **Rahmen der WS**¹³ zählend; und
- *dritte Zeile:* 42.4% Einheiten mit weniger als drei Beschäftigten und somit **ausserhalb des Rahmens der WS**.

Es wurde entschieden, die Modellierung auf die erste dieser beiden Subpopulationen (also die zum **Rahmen der WS** gehörenden Einheiten) zu konzentrieren, weil

- diese einen deutlich grösseren Anteil ausmacht als die zweite Subpopulation, wenn gewichtet nach Beschäftigung (16.0% vs. 0.4%) oder nach MWST-Umsätzen (15.0% vs. 4.3%) betrachtet; und
- die Zuverlässigkeit von Modellvorhersagen hier eher gewährleistet werden kann als für die zweite Subpopulation, da letztere aus Einheiten mit weniger als drei Beschäftigten besteht, die im Rahmen der WS nicht erfasst und somit auch in der Nettostichprobe nicht vertreten sind (Risiko der Verletzung der oben unter (c) genannten Bedingung).

Folglich erscheinen Modellvorhersagen für die Variable «Umsatz» machbar und sinnvoll für all jene Einheiten, die zum «ex post» konstruierten Rahmen der WS gehören, aber nicht in der Nettostichprobe erscheinen. Innerhalb dieser Subpopulation gibt es jedoch verschiedene Konstellationen des Vorhandenseins der potentiellen Hilfsvariablen aus STATENT und MWST. Es bleibt noch zu klären, wie mit diesen umzugehen ist. Hierzu werden für sämtliche der vier möglichen Konstellationen einerseits deren potentielle quantitative Bedeutung innerhalb der STAGRE beurteilt, andererseits aber auch die Anzahl der Beobachtungen, auf Grund derer sich das Modell schätzen lässt. Für letzteres wurde entschieden, sich nicht ausschliesslich auf Einheiten der STAGRE-Population abstützen, sondern die gesamte Grundgesamtheit der aktiven Unternehmen in der Schweiz zu berücksichtigen. Deren Anzahl Einheiten ist in Tabelle T2 dargestellt, gemäss dem gleichen

¹⁰ Da es während der hier relevanten Zeitspanne (2014 bis 2020) in den Grundgesamtheiten der Unternehmen und Unternehmensgruppen, aber auch im Stichprobenplan der WS keine einschneidenden Entwicklungen oder Änderungen gab, lässt sich davon ausgehen, dass sich die Erkenntnisse des Referenzjahres 2020 auch für die anderen Jahre Gültigkeit haben.

¹¹ Ausnahme: Ein kleiner Teil der Abteilung 86 wurde bis zum Referenzjahr 2015 durch die WS abgedeckt. Aus Kohärenzgründen werden in der STAGRE für diese Abteilung über die gesamte Zeitspanne keine Umsätze modelliert oder publiziert.

¹² Überdies enthalten für die erwähnten NOGA-Abteilungen die MWST-Daten mutmasslich nur lückenhafte Umsätze, da in diesen Abteilungen vornehmlich Leistungen erbracht werden, die von der MWST ausgenommen sind.

¹³ Vom «ex post» konstruierten Rahmen der WS ist hier die Rede, weil dieser auf Basis der aktuellsten verfügbaren Informationen zu Beschäftigtenzahl, NOGA-Zugehörigkeit und Rechtsform der Einheiten erstellt wurde. Zwischen diesen Informationen und jenen, die zum Zeitpunkt der Ziehung der Stichprobe der WS aus dem Betriebs- und Unternehmensregister (BUR) verfügbar waren, kann es Unterschiede geben. Diese sind jedoch geringfügig und für die wesentlichen Erkenntnisse der vorliegenden Analyse nicht bedeutsam. Der «ex post» konstruierte Rahmen wird hier gegenüber dem originalen WS-Rahmen bevorzugt, da zum Zweck der Modellierung ebenfalls auf die aktuellsten verfügbaren Informationen zurückgegriffen wird.

Schema wie in den drei Tabellen weiter oben. Mit diesem Vorgehen kann, dank einer grösseren Anzahl von Beobachtungen, mutmasslich die oben formulierte Bedingung (b) besser erfüllt werden, ohne dass dafür die Bedingung (c) in problematischer Art und Weise tangiert wäre.¹⁴

Für die verschiedenen **Konstellationen** des Vorhandenseins der potentiellen Hilfsvariablen zeigt sich folgendes:

- i. Einheiten, für die sowohl STATENT als auch MWST verfügbar sind (grün hinterlegt): In der STAGRE machen diese anzahlmässig 23.3% aus, etwas kleiner ist deren Bedeutung gewichtet nach Beschäftigung (15.4%) und Umsatz (15.0%). Eine Modellierung ist somit sinnvoll und dürfte, da beide Hilfsvariablen vorhanden sind, und da zur Schätzung des Modells 12 800 Beobachtungen zur Verfügung stehen (Tabelle T2), relativ zuverlässige Resultate liefern.
- ii. Einheiten mit Daten aus der STATENT, aber ohne MWST (gelb hinterlegt): Die Bedeutung solcher Einheiten ist in der STAGRE gering (anzahlmässig: 1.1%, gewichtet nach Beschäftigung: 0.6%).¹⁵ Eine Modellierung ist bedingt sinnvoll, und deren Resultate dürften im Vergleich zur Konstellation (i) weniger zuverlässig sein (da nur die Hilfsvariable Beschäftigung vorhanden ist, und auf Grund der geringeren für die Modellschätzung grundsätzlich zur Verfügung stehenden Zahl der Beobachtungen von 1 036),¹⁶ sollte aber dennoch in Erwägung gezogen werden.¹⁷

- iii. Einheiten mit Daten aus der MWST, aber ohne STATENT: Für diese Konstellation gibt es per Definition keine Einheiten, die in den «ex post» konstruierten Rahmen der WS fallen, da eine Anzahl von mindestens drei Beschäftigten gemäss STATENT ein Kriterium für die Zugehörigkeit zu diesem ist. Somit kann diese Konstellation hier ignoriert werden.

Als **Fazit** aus dieser Analyse wurde als Modellierungsstrategie festgelegt, **für die Konstellationen (i) und (ii) Modellvorhersagen für den Umsatz zu berechnen**. Zwar verbleiben so 42.4% der Einheiten, obschon sie eine für die WS relevante NOGA/Rechtsform-Kombination aufweisen, ohne verfügbare Umsätze (da weniger als drei Personen beschäftigend). Deren Beschäftigungssumme beträgt jedoch lediglich 0.4% und deren Umsätze gemäss MWST 4.3% des Totals aller STAGRE-Einheiten, was vertretbar erscheint angesichts der Tatsache, dass alle Einheiten, welche die Kriterien des WS-Rahmens erfüllen (also jene mit mindestens drei Beschäftigten), vollumfänglich abgedeckt werden.

Tabelle T2: Verteilung der Grundgesamtheit aller aktiven Unternehmen, Referenzjahr 2020

Anzahl der Einheiten...	Mit Daten verfügbar aus							
	STATENT und MWST		nur STATENT		nur MWST		(keine dieser beiden)	Total
... in der Nettostichprobe der WS	12'800	1.1%	1'036	0.1%	63	0.0%	35	13'934 1.2%
... ausserhalb Nettostichprobe, aber im Rahmen der WS	131'620	11.2%	15'200	1.3%	0	0.0%	0	146'820 12.5%
... in WS-relevanter NOGA/RF, aber mit < 3 Beschäftigten	126'499	10.8%	183'876	15.7%	96'812	8.3%	266'275	673'462 57.5%
... ausserhalb WS-relevanter NOGA/Rechtsform	20'569	1.8%	125'187	10.7%	16'980	1.5%	173'630	336'366 28.7%
Total	291'488	24.9%	325'299	27.8%	113'855	9.7%	439'940	1'170'582 100.0%

¹⁴ Es wird angenommen, dass sich Einheiten der STAGRE mit gegebenen Eigenschaften (Branche, Beschäftigungszahl, Rechtsform) in struktureller Hinsicht nicht wesentlich unterscheiden von ansonsten ähnlichen Einheiten, die keiner Unternehmensgruppe angeschlossen sind. Zudem besteht in der Spezifikation des Modells die Möglichkeit, gewisse Unterschiede mittels binärer Variablen für die Zugehörigkeit zu einer Unternehmensgruppe aufzufangen.

¹⁵ Hier handelt es sich im Wesentlichen um Einheiten, deren Umsatz entweder geringfügig ist (unterhalb der Schwelle von CHF 100 000, ab welcher die Deklaration der MWST obligatorisch ist) oder deren Leistungen von der Steuer ausgenommen sind. Zu den Einheiten ohne MWST-Umsätze werden hier auch jene gezählt, für die in den Unternehmensregisterdaten des BFS Imputationen

berechnet wurden, da diese Imputationen im vorliegenden Kontext nicht verwendet werden.

¹⁶ Für die Modellierung des Umsatzes dieser Einheiten wird ein Ansatz propagiert, der es erlaubt, auf Daten einer grösseren Schätzpopulation als lediglich der hier erwähnten 1 036 Einheiten zurückzugreifen. Die Details dazu werden in Kapitel 3 sowie im technischen Anhang erläutert.

¹⁷ Dies auch deshalb, weil für gewisse NOGA-Abteilungen in der STAGRE die Bedeutung dieser Einheiten deutlich grösser ist als gesamthaft über alle Branchen: Gewichtet nach Beschäftigung resultieren Anteile von 5% oder mehr in den Abteilungen 72, 85, 87 und 88.

2 Anforderungen und Spezifikation des verwendeten Modells

2.1 Ausgangslage zur Modellierung

Wie im einleitenden Kapitel begründet, wurde entschieden, den von der Wertschöpfungsstatistik (WS) erhobenen Umsatz als **Zielvariable** zu verwenden. Dies bedeutet, dass als massgebliche Komponenten für die Berechnung von aggregierten Umsätzen in der STAGRE zu verwenden sind:

- der gemäss WS erhobene Umsatz für jene Einheiten, für welche dieser vorhanden ist (also für die Einheiten der Nettostichprobe);
- ein Schätzwert des (hypothetischen) Umsatzes gemäss WS für jene Einheiten, die von der WS nicht erfasst wurden, modelliert auf der Basis verfügbarer Hilfsvariablen. Als Hilfsvariablen in Frage kommen: Umsatz gemäss MWST, Anzahl Beschäftigte und/oder vollzeitäquivalente Beschäftigung (VZÄ) gemäss STAGRE, Arbeitseinkommen gemäss AHV.

Die Konzeption eines statistischen Modells zur Schätzung des WS-Umsatzes war also zentraler Bestandteil unserer Vorgehensweise. Ein Modell kann die Realität niemals perfekt abbilden und muss sich zwangsläufig auf gewisse vereinfachende Hypothesen bezüglich der Zusammenhänge zwischen den verwendeten Variablen stützen. Deshalb war es unerlässlich, den Entscheid zugunsten einer schlussendlich zu verwendenden Modellvariante gestützt auf möglichst kohärente und nachvollziehbare Kriterien vorzunehmen. Diese Kriterien fassen auf eine empirische Evaluation der Modellresultate, ergänzt (wo sinnvoll) durch theoretische Überlegungen. Die schlussendlich verwendete Modellvariante sollte die folgenden **Anforderungen** möglichst gut erfüllen (die Reihenfolge dieser Auflistung gibt die Prioritätenfolge wieder, d.h. der ersten Anforderung wird die höchste Priorität eingeräumt):

- möglichst geringe Abweichung der Schätzwerte in Bezug auf Originalwerte der Nettostichprobe (auf Ebene der einzelnen Beobachtungen)
- möglichst geringe Abweichung von den Resultaten der direkten Schätzungen der WS (auf Ebene der Aggregate, für welche die direkten Schätzungen der WS konzipiert wurden)
- Vermeidung unplausibler modellbedingter Ausschläge im Zeitverlauf (auf Ebene der Aggregate, in denen die Aggregation der Modellresultate in der Statistikproduktion erfolgen wird)

Das schlussendlich verwendete Modell wird im unmittelbar folgenden Abschnitt beschrieben. Im danach anschliessenden Kapitel «Evaluation verschiedener Modellvarianten» wird erläutert, wie vorgegangen wurde, um die Eignung dieses verwendeten Modells gegenüber anderen theoretisch in Frage kommenden Varianten von Modellen zu begutachten und zu begründen; insbesondere unter Einbezug von konkreten Messgrössen, die die Erfüllung der soeben aufgelisteten Anforderungen erfassen.

2.2 Verwendetes Modell

Im verwendeten Modell ergibt sich der Umsatz gemäss Wertschöpfungsstatistik U_i eines Unternehmens i aus dem Umsatz gemäss MWST-Register T_i und den Arbeitseinkommen gemäss AHV L_i mittels folgendem Zusammenhang:

$$U_i = A_{g(i)} T_i^{\beta_{g(i)}} L_i^{\gamma_{g(i)}} e^{\delta' x_i} e^{\varepsilon_i}$$

Der Term $g(i)$ verweist auf die NOGA-Art (also die Branchenzugehörigkeit gemäss der detailliertesten verfügbaren Gliederungsstufe) des Unternehmens i . Der Einfluss, den die beiden verwendeten Hilfsvariablen T_i und L_i auf die Zielvariable U_i haben, variiert somit in Abhängigkeit der wirtschaftlichen Tätigkeit des Unternehmens. Der Vektor x_i enthält die folgenden erklärenden binären Variablen:

- ob i Teil einer multinationalen Unternehmensgruppe ist;
- ob i im Verlauf des jeweiligen Referenzjahres Waren exportiert hat;
- ob i im Verlauf des jeweiligen Referenzjahres Waren exportiert hat mit einem Wert, der den Wert des MWST-Umsatzes T_i übersteigt.

Der Anpassungsfaktor $A_{g(i)}$ trägt der Tatsache Rechnung, dass für gegebene Werte der erklärenden Variablen nicht zwangsläufig in jeder Branche der gleiche Wert der Zielvariablen resultieren muss. Der exponierte Zufallsterm ε_i absorbiert sämtliche in der Modellierung nicht berücksichtigten Faktoren. Aus der hier verwendeten multiplikativen Formulierung lässt sich, unter Anwendung von Logarithmen, der folgende lineare Zusammenhang herleiten:

$$\log U_i = \log A_{g(i)} + \beta_{g(i)} \log T_i + \gamma_{g(i)} \log L_i + \delta' x_i + \varepsilon_i$$

Für die Schätzung der Modellparameter (also der Vektoren $\log A_{g(i)}$, $\beta_{g(i)}$, $\gamma_{g(i)}$ und δ) wird diese lineare Formulierung verwendet. Die Schätzung wird wie folgt umgesetzt:

- Es wird, im Rahmen des «**Linear Mixed Model**»-Ansatzes (LMM), eine Kombination aus «Fixed»- und «Random»-Effekte verwendet. Die Parameter $A_{g(i)}$, $\beta_{g(i)}$ und $\gamma_{g(i)}$ sind als «Random»-Effekte spezifiziert, und zwar hierarchisch auf den Ebenen **noga2r** (siehe nachfolgenden Absatz) und NOGA-Art, wobei auch Korrelationen zwischen den drei Variablen modelliert werden. Die in δ erfassten Variablen werden als Fixeffekte spezifiziert.
- In der Nettostichprobe der WS wurden für einen nicht vernachlässigbaren Teil der Einheiten positive Umsätze erhoben, obschon diese Einheiten **keine MWST-Umsätze deklariert** haben (mehr als fünf Prozent, siehe die Erläuterungen des vorhergehenden Kapitels sowie Tabelle T2 für das Beispiel des Referenzjahres 2020). Um eine Schätzung des Umsatzes auch für Einheiten mit fehlendem MWST-Umsatz (aber mit vorhandener Lohnsumme gemäss AHV) zu ermöglichen, wird im Rahmen eines «hybrid separat/gepoolten» Ansatzes (siehe Technischer Anhang) zusätzlich eine Variante der Modellgleichung unter Weglassung der erklärenden Variablen $\log T_i$ geschätzt.
- Ein logarithmiertes Modell liefert eine Vorhersage für $E(\log y)$, dem Erwartungswert des Logarithmus der abhängigen Variable. Der Erwartungswert $E(y)$ – also der untransformierten abhängigen Variable – entspricht jedoch nicht $\exp[E(\log y)]$, da der Logarithmus keine lineare Transformation ist. Dies bedeutet, dass für eine unverzerrte Vorhersage des untransformierten Umsatzes eine **Skalierung** erforderlich ist. Zu diesem Zweck wurde ein jeweils separat innerhalb jeder **noga2r**-Branche differenzierter Skalierungsfaktor $\exp[E(\hat{\varepsilon}^2)/2]$ verwendet, wobei Extremwerte in $\hat{\varepsilon}$ (dem Vektor der vorhergesagten Modellresiduen) ignoriert wurden. Als extrem eingestuft wurden hier Beobachtungen, deren Absolutwert für $\hat{\varepsilon}$ die die Schwelle von drei geschätzten Standardabweichungen überschritten.
- Die Schätzung der Modelle und die Vorhersage der Umsätze erfolgt **isoliert** für jedes der abgedeckten Referenzjahre 2014 bis 2020.

Im Kontext der Modellkonstruktion und -evaluation wird verschiedentlich auf die Branchenklassifizierung **noga2r** zurückgegriffen. Diese stützt sich auf die NOGA-Abteilungen (erste zwei Stellen der NOGA-Klassifizierung), wobei – auf Grund der geringen Anzahl der darin enthaltenen Einheiten und in Übereinstimmung mit dem Robustifizierungsverfahren der Hochrechnung der WS – die folgenden vier Abteilungsgruppen jeweils

zusammengefasst wurden: 05–09, 11–12, 19–20, 38–39.¹⁸ Nach dieser Zusammenlegungen sind, über alle der hier untersuchten Referenzjahre hinweg, sämtliche Ausprägung mit einer Zahl von mindestens 27 auswertbaren Beobachtungen belegt.

Im nächsten Kapitel firmiert das bis hierhin beschriebene Modell unter dem Namen **dms_r10_s23**. Das schlussendlich verwendete Modell **dms_r10_46cs_0cap_i_s23** wurde von diesem abgeleitet und unterscheidet sich davon durch die drei folgenden Verfeinerungen (die im Abschnitt 3.4 detailliert erläutert werden):

- In einigen NOGA-Arten resultieren negative und somit unplausible Koeffizienten-Schätzwerte für $\beta_{g(i)}$ oder $\gamma_{g(i)}$. Für die betroffenen NOGA-Arten wird automatisiert auf «Ausweichmodelle» zurückgegriffen, die auf den Einbezug der betroffenen Hilfsvariablen verzichten (detaillierte Ausführungen zu diesem Vorgehen finden sich im Technischen Anhang).
- Für die NOGA-Abteilung 46 (Grosshandel) wurde ein linear homogener Zusammenhang zwischen den Hilfsvariablen und der abhängigen Variablen (also das Erfordernis, dass $\beta_{g(i)} + \gamma_{g(i)} = 1$) erzwungen.
- Nach punktuellen Analysen von Fällen mit übermässig hohen vorhergesagten Umsätzen wurde bei sechs Einheiten erzwungen, für eines oder mehrere Referenzjahre auf Schätzwerte eines anderen (vereinfachten) Modells zurückzugreifen, da die Vorhersagen gemäss Hauptmodell als unplausibel hoch zu erachten waren.

¹⁸ Siehe «Methodenbericht Wertschöpfungsstatistik – Revision 2009: Statistische Datenaufbereitung und Hochrechnung», BFS, Neuchâtel, 2014

3 Evaluation der verschiedenen Modellvarianten

3.1 Grundlegende Herausforderungen

Bevor vertieft auf die Kriterien und Vorgehensweise der Modellevaluation eingegangen wird, ist es sinnvoll, zunächst die Aufmerksamkeit auf zwei grundlegenden Herausforderungen zu lenken, die sich bei der Evaluation von modellbasierten Schätzungen für die Variable Umsatz stellen: Erstens die extrem rechtsschiefe Verteilung dieser Variablen; und zweitens die Auswirkungen des Einbezugs zusätzlicher und potentiell aktuellerer (in Bezug auf den Rahmen der Wertschöpfungsstatistik) Daten.

Charakteristisch für die Variable Umsatz und erschwerend für die Konzeption und Evaluation von Modellen ist, dass diese Kennzahl **extrem ungleich (also rechtsschief)** verteilt ist. Diese ungleiche Verteilung manifestiert sich auf Ebene sowohl der einzelnen Einheiten als auch der für die Statistik relevanten Aggregaten nach Wirtschaftstätigkeit.¹⁹ Für eine möglichst gute Erfüllung der im vorhergehenden Kapitel erläuterten Anforderungen ist es deshalb wichtig, Kriterien zu definieren, die nicht in einseitiger Art und Weise durch die Modellresultate bezüglich dieser sehr umsatzstarken Einheiten oder Branchen getrieben sind, sondern auch berücksichtigen, wie gut ein Modell für Einheiten und Branchen mit geringen oder mittleren Umsätzen abschneidet.

Der Einbezug in die Modellschätzungen von **Daten, die zum Zeitpunkt der Konstruktion des Rahmens der Wertschöpfungsstatistik noch nicht verfügbar waren**, hat potentiell Auswirkungen darauf, inwiefern sich die für die Modellevaluation spezifizierten Anforderungen aussagekräftig messen lassen. Als Konsequenz daraus wird für die Evaluation in zwei Etappen vorgegangen:

- In einer **ersten Etappe** werden Modelle geschätzt und evaluiert, die sich ausschliesslich auf Informationen der WS abstützen; insbesondere auf die erklärenden Variablen **noga2r** (siehe vorhergehendes Kapitel) und Beschäftigung in Vollzeitäquivalenten (VZÄ)²⁰ aus dem Rahmen, sowie auf die Zielvariable Umsatz aus der

Erhebung. Zusätzliche Datenquellen, die potentiell präzisere und/oder aktuellere Schätzungen des Umsatzes erlauben (wie STATENT, MWST-Umsätze), werden also vorerst ignoriert.²¹ Der Zweck dieser ersten Etappe besteht darin, verschiedene **«Basisvarianten» der Modellspezifikationen** auf ihre Eignung zu überprüfen, indem die aus diesen Spezifikationen resultierenden Modellvorhersagen mit den direkten Schätzungen der WS verglichen werden (und somit eine möglichst aussagekräftige Einschätzung bezüglich der in Abschnitt 2.1 erläuterten Anforderung 2 zu erhalten). Die Wahl der Basisvariante beinhaltet jene Entscheidungen, die nicht den spezifischen Einbezug der zusätzlichen Datenquellen betreffen, wie etwa die funktionale Form (logarithmisch oder in absoluten Werten) oder der Einbezug von Gewichtung und/oder Bereinigung um Ausreisser.

- Ist die Entscheidung für eine (oder mehrere, als annähernd gleichwertig eingestufte) Basisvariante getroffen, folgt darauf aufbauend in einer **zweiten Etappe** die Untersuchung und Evaluation der **«Detailvarianten» der Modellspezifikationen**. Diese Detailvarianten betreffen insbesondere die Auswahl der zu verwendenden Hilfsvariablen aus den in der ersten Etappe ignorierten zusätzlichen Datenquellen und deren konkrete Formulierung.

Wenn auf diese Etappierung verzichtet würde, d.h. wenn von Beginn weg lediglich Modellschätzungen analysiert würden, die auch Daten ausserhalb des Rahmens und der Erhebung der WS berücksichtigen, wäre die Modellevaluation dadurch erschwert, dass der Einbezug dieser Daten die Übereinstimmung zwischen Modellvorhersage und direkten Schätzungen (Anforderung 2) potentiell verringert. Der Stand der Daten zum Zeitpunkt der Erstellung des Rahmens der WS (die für die Berechnung der direkten Schätzungen verwendet werden) kann sich in Inhalt und Umfang erheblich von den aktuellsten Daten (wie sie für die Modellschätzungen verwendet werden sollen) abweichen.²² In der Folge

¹⁹ In der Nettostichprobe der WS erzielt das umsatzstärkste Promille der Einheiten mehr als ein Drittel und das umsatzstärkste Prozent der Einheiten zwei Drittel des totalen erhobenen (ungewichteten) Umsatzes («gepoolte» Betrachtung über die Referenzjahre 2014-2020).

Der gesamte Umsatz (gemäss Direktschätzung) der NOGA-Abteilung 46 (Grosshandel) beträgt in allen Referenzjahren über CHF 1000 Mia., während in allen anderen NOGA-Abteilungen das Total stets unter oder nur knapp über CHF 100 Mia liegt.

²⁰ Für die Konstruktion der Schichten im Stichprobenplan der WS werden nicht VZÄ, sondern die Beschäftigungsvariable **betot** gemäss BUR verwendet; ausserdem dient **betot** als Hilfsvariable im bis zum Referenzjahr 2020 verwendeten

Imputationsmodell für fehlende MWST-Umsätze. Hingegen sind die VZÄ des WS-Rahmens von zentraler Bedeutung für die Konstruktion der Hochrechnungsgewichte der WS, weshalb davon ausgegangen wird, dass die VZÄ (und nicht **betot**) die relevanteste Hilfsvariable im Kontext der ersten Etappe darstellt.

²¹ Im Zuge einer Revision der WS, die ab dem Referenzjahr 2021 umgesetzt wird, werden in Zukunft zusätzlich zur Beschäftigung auch die MWST-Umsätze als Hilfsvariable in den Stichprobenplan und die Gewichtung einfließen.

²² Dies äussert sich etwa darin, dass aggregiert auf Ebene **noga2r** in mehr als der Hälfte der Fälle die vollzeitäquivalente Beschäftigung um mehr als 2.5% abweicht, wenn statt des ursprünglichen WS-Rahmens die aktuellsten Daten herangezogen werden. Es handelt sich dabei um den kombinierten Effekt von

können – selbst wenn ein bezüglich der Anforderung 1 optimales Modell gewählt wurde – auch auf Ebene der relevanten Aggregate die Resultate der Modellschätzungen substanziell von jenen der direkten Schätzungen abweichen. Lässt sich jedoch – dank der hier vorgenommenen Etappierung – zeigen, dass allenfalls problematisch gross erscheinende Abweichungen eine Folge der Datenlage (und nicht der Modellspezifikation) sind, kann dies die Wahl eines bezüglich der Anforderung 1 optimalen Modells stützen.

3.2 Kriterien und Grundsätze

Zur Evaluation der Modelle dienen in erster Linie die zwei folgenden Arten von Kriterien:

- Ein globales Qualitätsmass (**QM-Global**) beschreibt die Präzision eines Modells auf Ebene der einzelnen Beobachtungen (Unternehmen) und misst somit die Qualität in Bezug auf die in Kapitel 2.1 erwähnte **Anforderung 1**. Es erlaubt – insbesondere, wenn es nicht unter Verwendung einer geeigneten Gewichtung implementiert wird – nur begrenzt Rückschlüsse zur Prognosegüte ausserhalb der Nettostichprobe, da Referenzwerte auf Ebene der einzelnen Beobachtungen nur innerhalb der Nettostichprobe zur Verfügung stehen.
- Ein Qualitätsmass auf aggregierter Ebene (**QM-Aggregat**) beschreibt die Präzision eines Modells bezüglich der Resultate auf aggregierter Ebene (z.B. nach Branchen). Es lässt sich – unter Verwendung der direkten Schätzungen der WS als Referenzwerte – bezogen auf den gesamten Rahmen der WS berechnen, misst somit die Qualität in Bezug auf die **Anforderung 2** und erlaubt zu einem gewissen Grad Rückschlüsse zur Prognosegüte ausserhalb der Nettostichprobe.

Wie bereits angetönt, erhält QM-Aggregat in der ersten Etappe Aufmerksamkeit, da Anforderung 2 im Wesentlichen in der ersten Etappe untersucht wird. Hingegen ist QM-Global das relevantere Kriterium in der zweiten Etappe, da dieses Mass Hinweise dazu liefert, inwiefern der Einbezug von Daten ausserhalb Rahmen und Erhebung der WS eine präzisere Schätzung des Umsatzes ermöglicht (und somit Anforderung 1 erfüllt wird), und zwar auch dann, wenn als Folge des Einbezugs dieser Daten die Übereinstimmung der Modellresultate mit den direkten Schätzungen der WS abnimmt. **Anforderung 3** wird schliesslich mittels grafischer Darstellungen und einer vertieften Analyse von einzelnen einflussreichen Beobachtungen überprüft und, wo nötig, mit punktuellen Korrekturen gewährleistet. Die Details und die konkrete Umsetzung werden im Abschnitt 3.4 dargelegt.

QM-Global

Für QM-Global würden zunächst die vom Modell unmittelbar gelieferten Standardkennzahlen (wie R^2 , residualer Standardfehler, p-Werte der Koeffizienten und des Gesamtmodells) in Frage kommen. Es gilt jedoch zu beachten, dass zwischen Modellen mit fundamental unterschiedlicher Spezifikation (z.B. logarithmiert vs. in absoluten Werten) diese Kennzahlen nur bedingt oder gar nicht miteinander vergleichbar sind. Eine bessere Vergleichbarkeit lässt sich erzielen, indem der Vektor (über alle Unternehmen der Nettostichprobe) der geschätzten Werte der Zielvariablen berechnet²³ und auf diesen geeignete Beurteilungskriterien angewendet werden. Konkret dient dazu der **«Root Mean Squared Error» RMSE** (Definition siehe Kasten), der ein Mass für den mittleren Fehler der Vorhersage darstellt, und für den somit möglichst tiefe Werte anzustreben sind (ein Wert von null entspricht einer perfekten Vorhersage). Wie sich weiter unten zeigen wird, beeinträchtigt die eingangs erwähnte ausgeprägt rechtsschiefe Verteilung des Umsatzes die Interpretierbarkeit dieses Masses. Deshalb werden, nebst den «rohen» RMSE-Werten, auch zwei Varianten von RMSE

Root Mean Squared Error (RMSE)

$$RMSE_w = \sqrt{\frac{1}{\sum_{i=1}^N w_i} \sum_{i=1}^N (\hat{y}_i - y_i^R)^2 w_i}$$

Der **Root Mean Squared Error (RMSE)** aggregiert die Abweichungen zwischen einem vom Modell geschätzten Vektor $\hat{y}_i$ und einem Vektor von Referenzwerten y_i^R (effektiv beobachtete Werte, oder Werte einer Referenzschätzung) einer Variablen (hier der Umsatz gemäss WS), unter Einbezug eines geeigneten Gewichts w_i (wird auf eine Gewichtung verzichtet, ist im obigen Ausdruck einfach $w_i = 1$ zu setzen). Die zugrundeliegende Beobachtungsstufe i kann auf einzelne Einheiten (für ein globales Mass) oder auf bereits aggregierte Werte (z.B. nach Branchen) verweisen. Für den RMSE ist ein möglichst tiefer Wert wünschenswert: Kleinere Werte bedeuten im Mittel kleinere Schätzfehler und somit eine präzisere Modellvorhersage. Zum Zweck einer die gesamte Zeitspanne abdeckenden Evaluation der Schätzwerte wird die Summierung nicht über die Einheiten eines einzelnen Referenzjahres, sondern über alle Referenzjahre t und Einheiten i gemäss der folgenden Gleichung vorgenommen:

$$RMSE_w = \sqrt{\frac{1}{\sum_{t=1}^T \sum_{i=1}^{N_t} w_{it}} \sum_{t=1}^T \sum_{i=1}^{N_t} (\hat{y}_{it} - y_{it}^R)^2 w_{it}}$$

Änderungen in der NOGA-Zugehörigkeit, der Beschäftigung und des Aktivitätsstatus von Einheiten.

²³ Im Falle eines logarithmisch spezifizierten Modells (wo $y = \log U$) bedeutet dies, dass die Schätzwerte mittels e^y in absolute Werte für U zurücktransformiert und allenfalls skaliert werden müssen (siehe weiter unten).

herangezogen, in welchen der Einfluss der umsatzstarken Branchen und Beobachtungen mittels geeigneter Gewichtung eingegrenzt wird. Es kommen die folgenden Gewichte zur Anwendung:

- **HR:** Dies ist das von der WS für die Berechnung der direkten Schätzungen verwendete Hochrechnungsgewicht. Somit erhalten Schätzfehler von Beobachtungen, die für die Hochrechnung der WS bedeutsam sind (also jene, deren erhobene Werte als Folge von Stichprobenplan, Rücklaufquote und Robustifizierung für eine wesentliche Zahl von nicht erhobenen Einheiten als repräsentativ behandelt werden) ein höheres Gewicht.
- **HR & Branchen** (mit Branchenfaktoren): Dies ist das soeben erwähnte Hochrechnungsgewicht der WS, dividiert durch einen nach **noga2r** differenzierten Faktor, der den Gesamtumsatz einer Branche im Verhältnis zum durchschnittlichen Gesamtumsatz aller Branchen ausdrückt (als Gesamtumsätze dienen hier die hochgerechneten Resultate der WS). Der Effekt dieses Faktors lässt sich mit dem folgenden Beispiel veranschaulichen: Angenommen, die Vorhersage des Modells für den Umsatz einer Einheit weicht um eine Million Franken vom erhobenen Wert ab. Gehört die Einheit einer Branche mit einem Gesamtumsatz von 100 Milliarden Franken an, so ist der Effekt dieser Abweichung auf das so gewichtete Qualitätsmass zehn Mal geringer, als wenn die Einheit einer Branche mit einem Gesamtumsatz von 10 Milliarden Franken angehören würde (gegeben, dass die Einheit in beiden Fällen das gleiche Hochrechnungsgewicht hat). Somit werden Schätzfehler in Branchen mit einem höheren Gesamtumsatz heruntergewichtet, während Schätzfehler in Branchen mit kleinem Gesamtumsatz ein höheres Gewicht erhalten.

Ergänzend zu diesen RMSE-Varianten werden wo sinnvoll Residuenplots erstellt, die insbesondere auch Hinweise liefern können, ob die getroffene Modellspezifikation Schwächen in Bezug auf Linearitäts- oder Homoskedastizitätsannahmen aufweist.

QM-Aggregat

Ausgangspunkt ist hier der Vektor der geschätzten Werte von Interesse (also des WS-Umsatzes), nun jedoch für alle Unternehmen des Rahmens der WS (und nicht wie bei QM-Global nur für die Nettostichprobe). Da bei QM-Aggregat der Fokus auf den Prognoseeigenschaften ausserhalb der Nettostichprobe liegt, werden für die Einheiten der Nettostichprobe an Stelle der geschätzten Werte die erhobenen Werte eingesetzt.²⁴ Die Werte des daraus resultierenden Vektors (im Folgenden der Einfachheit halber als «erhebungsergänzte Modellvorhersagen», kurz EMV bezeichnet) werden nach der Branchenklassifizierung **noga2r**

(siehe oben) aggregiert, und die Resultate dieser Aggregation lassen sich nun mit den auf gleicher Aggregationsstufe vorliegenden direkten Schätzungen der WS vergleichen. Zu beachten ist hier, dass es sich bei Letzteren ebenfalls um Schätzwerte handelt. In der Folge sind Abweichungen zwischen den EMV-Aggregaten und den direkten Schätzungen nicht zwingend auf Unzulänglichkeiten des Modells zurückzuführen; vielmehr können diesen Abweichungen Schätzfehler sowohl der direkten Schätzungen als auch des Modells zugrunde liegen.

Um trotz dieser Schwierigkeiten ein Qualitätsmass zu berechnen, das eine sinnvolle Interpretationsmöglichkeit bietet, werden die Abweichungen zwischen EMV-Aggregaten $\hat{y}_{jt}^{EMV}$ und direkten Schätzungen $\hat{y}_{jt}^{WS}$ so über die Referenzjahre $t = 1..T$ und **noga2r**-Branchen $j = 1..M$ summiert, dass Abweichungen in Branchen mit einer hohen Präzision der direkten Schätzungen höher gewichtet werden als solche in Branchen mit weniger präzisen direkten Schätzungen. Hierzu wird zunächst für jede Beobachtung (Branche x Referenzjahr) der **z-Wert** berechnet, also die standardisierte (durch die geschätzte Standardabweichung der direkten Schätzung $\hat{\sigma}_{\hat{y}_{jt}^{WS}}$ geteilte) Abweichung:

$$z_{jt} = \frac{\hat{y}_{jt}^{EMV} - \hat{y}_{jt}^{WS}}{\hat{\sigma}_{\hat{y}_{jt}^{WS}}}$$

Ausgehend von diesen z-Werten werden zwei Varianten von Qualitätsmassen berechnet und betrachtet:

- der **Mittelwert** dieser z-Werte:

$$\frac{1}{MT} \sum_{t=1}^T \sum_{j=1}^M z_{jt}$$

- die Quadratwurzel des Mittelwerts der quadrierten z-Werte (**«Root Mean Square»** oder **RMS**).²⁵

$$\sqrt{\frac{1}{MT} \sum_{t=1}^T \sum_{j=1}^M z_{jt}^2}$$

Die beiden Varianten ermöglichen dabei eine differenzierte Einschätzung darüber,

- ob die EMV-Aggregate insgesamt systematisch höher oder tiefer ausfallen als die direkten Schätzungen (anhand des **Mittelwerts** der z-Werte, der idealerweise möglichst nahe bei null liegen sollte); und
- wie stark die Abweichungen zwischen EMV-Aggregaten und den direkten Schätzungen gestreut sind (anhand

²⁴ Dies ist kohärent mit der Vorgehensweise zur Berechnung der letztendlich für die Publikation vorgesehenen Aggregate, wo erhobene Werte, falls vorhanden, ebenfalls eingesetzt werden.

²⁵ Konzeptionell ähnelt dieses Mass dem bereits erwähnten RMSE, mit dem Unterschied, dass die Summierung das Quadrat nicht der «rohen», sondern der standardisierten Abweichungen verwendet.

des **RMS** der z-Werte; hier sind ebenfalls möglichst tiefe Werte wünschenswert, wobei Werte bis eins oder leicht darüber grundsätzlich als unproblematisch angesehen werden können²⁶).

Grundsätze zur Konzeption und Evaluation der Modelle

Für die Berechnung der Modelle werden die Daten der zur Verfügung stehenden Referenzjahre 2014 bis 2020 separat betrachtet; das heisst, die Modelle werden für jedes Referenzjahr jeweils isoliert geschätzt. Auf eine über Referenzjahre «gepoolte» Schätzung der Modelle wird also verzichtet.²⁷ Hingegen wird für die Evaluation der Modelle grundsätzlich die gesamte Zeitspanne zusammengefasst betrachtet; die Berechnung der Qualitätsmasse beinhaltet also eine Summierung über alle Referenzjahre. Nach diesen Grundsätzen wird sowohl für die erste (Basisvarianten) als auch für die zweite (Detailvarianten) Etappe vorgegangen. In der zweiten Etappe finden sich zusätzlich auch nach Branchen (NOGA-Abschnitte) und Referenzjahren aufgeschlüsselte Evaluationsresultate der Modelle mit dem Zweck, allfällige Unterschiede in der Beschaffenheit der Daten zu identifizieren, die für die Modellwahl relevant sein könnten.

In sämtlichen Modellvarianten wird Heterogenität der Effekte in Bezug auf die Branchengliederung **noga2r** (siehe Abschnitt 2.2) unterstellt; also derjenigen Gliederungsstufe, die auch für die Aggregation der Modellresultate zur Bildung von QM-Aggregat zum Einsatz kommt. Dies geschieht, indem die Modellkonstante (wo spezifiziert) sowie die metrische(n) erklärende(n) Variable(n) mit Fix- oder Random-Effekten interagiert werden. Die Option, darüber hinaus detailliertere NOGA-Gliederungen zu berücksichtigen, wird in der zweiten Etappe untersucht und beschrieben.

3.3 Erste Etappe (Basisvarianten)

Grundsätzliche Überlegungen

Die Evaluation der Basisvarianten konzentriert sich – wie weiter oben angetönt – auf diejenigen Modelleigenschaften, die sich unabhängig vom Einbezug der zusätzlichen Datenquellen (ausserhalb der WS) stellen und evaluieren lassen, insbesondere:

- Formulierung in **Logarithmen vs. in absoluten Werten** sowie Wahl der **Skalierung** der Schätzergebnisse in logarithmischen Modellen
- **Erzwungene lineare Homogenität oder nicht** der abhängigen Variablen (Umsatz) in Bezug auf die metrischen Hilfsvariablen²⁸
- **Bereinigung um Ausreisser**
- **Gewichtung**

Der vorliegende Abschnitt präzisiert die mit diesen Modelleigenschaften verbundenen Überlegungen und Herausforderungen, und erläutert die sich daraus ergebenden, für die Modellevaluation zu berücksichtigenden Varianten. Eine schematische Übersicht über diese Basismodellvarianten findet sich in Tabelle T3. Als Hilfsvariable dient jeweils (wie weiter oben erläutert) die vollzeitäquivalente Beschäftigung gemäss STATENT.

Logarithmen vs. in absoluten Werten

A priori sind zur Modellierung der Zielvariablen Spezifikationen in absoluten Werten ($y = \alpha x_1 + \beta x_2 + \dots + \varepsilon$) wie auch in logarithmierten Werten ($\log y = \alpha + \beta \log x_1 + \gamma \log x_2 + \dots + \varepsilon$) denkbar.²⁹ Eine logarithmische Variante bietet dabei die folgenden Vorteile:

- Probleme auf Grund bedingter Heteroskedastizität (bezüglich der metrischen erklärenden Variablen) und stark von der Normalität abweichender Verteilung der Residuen dürften weniger gravierend sein (dies lässt sich mit geeigneten Residuenplots verifizieren).
- Ein logarithmisches Modell lässt Abweichungen von linearer Homogenität des Umsatzes in Bezug auf die Hilfsvariablen (siehe weiter unten) zu, indem es eine konstante Elastizität zwischen den Hilfsvariablen und

²⁶ Letzteres auf Grund der folgenden Überlegung: Die hypothetischen z-Werte der direkten Schätzungen in Bezug auf die (nicht beobachtbaren) tatsächlichen Populationswerte sind annähernd standardnormal verteilt (unter der Annahme, dass die Standardabweichungen der direkten Schätzungen konsistent geschätzt sind); folglich beträgt der Erwartungswert des Quadrats dieser z-Werte eins (da diese einer χ^2 -Verteilung mit Freiheitsgrad eins folgen) und somit tendiert auch deren RMS gegen eins. Da es keinen Sinn ergäbe, von den EMV-Aggregaten zu erwarten, dass sie näher bei den direkten Schätzungen zu liegen kommen als die tatsächlichen Populationswerte, ist ein Wert von eins für den RMS der z-Werte zwischen EMV-Aggregaten und direkten Schätzungen vollkommen akzeptabel.

²⁷ Mit einer gepoolten Vorgehensweise sind theoretisch Effizienzsteigerungen denkbar; dies jedoch zum Preis einer erhöhten Modellkomplexität und dadurch bedingt grösserem Rechenaufwand. Zudem würden Modellvorhersagen gestützt

auf eine gepoolte Vorgehensweise das Erfordernis einer sorgfältig konzipierten Revisionsstrategie mit sich bringen (da bei jedem Hinzufügen eines neuen Referenzjahres die Modellvorhersagen für sämtliche bisherigen Referenzjahre von Änderungen betroffen wären).

²⁸ Dies ist im Wesentlichen für logarithmische formulierte Modelle relevant, da – wenn keine expliziten Restriktionen spezifiziert werden – lineare Modelle mit abhängigen und erklärenden Variablen in jeweils logarithmischer Formulierung einen exponentiellen (und somit im allgemeinen Fall nicht linear homogenen) Zusammenhang unterstellen.

²⁹ Innerhalb der ersten Etappe dient als Hilfsvariable x_1 ausschliesslich die VZÄ-Beschäftigung gemäss Rahmen der WS, somit gibt es hier kein x_2 (erst in der zweiten Etappe ist vorgesehen, mehr als eine Hilfsvariable zu verwenden).

der Zielvariablen unterstellt, welche a priori nicht eins betragen muss (siehe technischer Anhang).³⁰

- Falls «Linear Mixed Models» (siehe weiter unten, zweite Etappe) geschätzt werden sollen, ist eine logarithmische Formulierung geeigneter. Dies, da diese Modelle eine Normalverteilung der zu schätzenden «Random Effects» unterstellen. Allfällige Abweichungen von dieser Verteilungsannahme dürften weniger gravierend sein, wenn das Modell logarithmiert spezifiziert wird.

Nachteile der logarithmischen Variante sind:

- Einheiten, die in einer der metrischen Variablen einen Wert von null aufweisen, können nicht ohne weiteres in der Modellschätzung und der Modellvorhersage berücksichtigt werden, obschon sie potentiell relevante Informationen enthalten.³¹
- Ein logarithmiertes Modell liefert eine Vorhersage für $E(\log y)$, dem Erwartungswert des Logarithmus der abhängigen Variable. Letztendlich von Interesse ist allerdings der Erwartungswert der untransformierten abhängigen Variable, $E(y)$. Dieser ist nicht gleich $\exp[E(\log y)]$, da der Logarithmus keine lineare Transformation darstellt. Im hypothetischen Fall einer normalen Verteilung der Residuen ε im logarithmierten Modell mit Standardabweichung σ gilt etwa $E(y) = \exp[E(\log y) + \sigma^2/2]$. Dies bedeutet, dass – Normalverteilung vorausgesetzt – zur Vorhersage des (untransformierten) Umsatzes eine **Skalierung** um den Faktor $\exp(\sigma^2/2)$ erforderlich ist. Die Implementierung einer geeigneten Skalierung kann aus diversen Gründen (Unsicherheit bezüglich der Schätzung von σ^2 , nichtnormale Verteilung der Residuen, bedingte Heteroskedastizität) Probleme aufwerfen und erfordert deshalb eine sorgfältige Evaluation.

Für die Evaluation werden folgende Varianten berücksichtigt:

- *Logarithmisch: **bm...**, mit Implementierung der Skalierung wie folgt:*
 - o **bm**: ohne Skalierung
 - o **bm_s0**: homogene Skalierung, Verwendung von $\exp(\sigma^2/2)$ als Skalierungsfaktor, ohne Bereinigung von Extremwerten
 - o **bm_s1**: heterogene (nach **noga2r**) Skalierung, Verwendung von $\exp[E(\varepsilon^2)/2]$ als Skalierungsfaktor, ohne Bereinigung von Extremwerten
 - o **bm_s23**, **bm_s23a**, **bm_s24**, **bm_s24a**: heterogene (nach **noga2r**) Skalierung, Verwendung

von $\exp[E(\varepsilon^2)/2]$ als Skalierungsfaktor, diverse Varianten der Bereinigung von Extremwerten in ε (siehe folgenden Abschnitt)

- Absolut: **bm_abs...** (diverse Varianten, weiter unten erläutert)

Erzwungene lineare Homogenität

A priori gibt es für die Modellierung valable Argumente sowohl zugunsten eines linear homogenen Zusammenhangs zwischen den metrischen Hilfsvariablen und der abhängigen Variablen (WS-Umsatz), als auch zugunsten einer Abweichung von einem solchen Zusammenhang. Letztere ist insbesondere ein Merkmal von logarithmischen Modellen der Form $\log(y) = a + \beta \log(x) + \varepsilon$ (falls nicht, mittels der Einschränkung $\beta = 1$, lineare Homogenität erzwungen wird) sowie von Modellen in absoluten Werten unter Einbezug einer Regressionskonstante. Da eine Abweichung von einem linear homogenen Zusammenhang, falls diese nicht allzu stark ausfällt (also β nahe eins ist im Falle logarithmischer Modelle), ökonomisch plausibel sein kann, werden beide Eventualitäten in die Evaluation einbezogen, allerdings nur für die logarithmische Variante.³²

Für die Evaluation wird folgende Variante berücksichtigt (nur im Rahmen logarithmischer Modelle):

- **bm_cs_s23**: mit erzwungener linearer Homogenität
- (Referenzvariante dazu, ohne erzwungene lineare Homogenität: **bm_s23**)

Bereinigung um Ausreisser

Statistische Ausreisser (Beobachtungen mit Extremwerten in den Residuen des unterstellten Modells) können zu übermässiger Variabilität der Schätzkoeffizienten (und letztlich der Schätzwerte) führen. Dieses Problem lässt sich mildern, indem versucht wird, solche Ausreisser in den Daten zu identifizieren und in den letztlich zu verwendenden Schätzungen entweder auszuschliessen oder (wie im nachfolgenden Unterabschnitt diskutiert) herunterzugewichten. Da jeder Ausschluss von Beobachtungen das Risiko birgt, relevante Informationen irrtümlicherweise zu ignorieren, werden Varianten mit und ohne eine solche Ausreisserbereinigung evaluiert. Die hier evaluierten Varianten der Ausreisserbereinigung funktionieren nach dem Grundsatz, dass zunächst die Residuen aus der Referenzvariante (unbereinigte Schätzung) ermittelt werden, und dass alle Beobachtungen, deren Residuum in ein bestimmtes Quantil von extremsten Werten fällt, von der

³⁰ Die Bedingung einer linear homogenen Beziehung lässt sich, falls erwünscht, problemlos auch in einem logarithmischen Modell unterbringen.

³¹ Total über die Referenzjahre 2014-2020 wurde bei 77 Beobachtungen ein positiver WS-Umsatz, aber ein Wert von null für den MWST-Umsatz erhoben; dies entspricht weniger als 0,1% aller Beobachtungen mit positivem WS-Umsatz. Gemessen an deren hochgerechnetem Umsatz in der WS ist das Gewicht dieser Beobachtungen nochmals geringer. Der Einfachheit halber werden Beobachtungen mit einem Wert von null für den MWST-Umsatz in der Modellierung gleich

behandelt wie Beobachtungen ohne vorhandenen MWST-Umsatz (siehe Technischer Anhang). Eine analoge Betrachtung bezogen auf die Variable Arbeitseinkommen zeigt, dass dieses Phänomen hier deutlich weniger ausgeprägt ist (eine betroffene Beobachtung).

³² Auf eine Evaluation von Modellen in absoluten Werten unter Einbezug einer Regressionskonstante wird verzichtet, da solche Modelle Probleme in der Form negativer Schätzwerte verursachen können; alle der hier evaluierten Modelle in absoluten Werten unterstellen somit lineare Homogenität.

Schätzpopulation ausgeschlossen und das Modell danach erneut geschätzt wird, um so die bereinigte Variante zu erhalten.³³

Für die Evaluation werden folgende Varianten berücksichtigt:

- *Logarithmisch:*
 - o **bm_oc_s23:** wiederholte Schätzung des Modells, nach Ausschluss jener Beobachtungen, deren Absolutwert des Residuums dem Perzentil mit den grössten Schätzfehlern angehört
 - o (Referenzvariante dazu: **bm_s23**)
- *Absolut:*
 - o **bm_abs_ew_oc:** wiederholte Schätzung des Modells, nach Ausschluss jener Beobachtungen, deren Absolutwert des Residuums dem Perzentil mit den grössten Schätzfehlern angehört
 - o (Referenzvariante dazu: **bm_abs_ew**; siehe untenstehenden Abschnitt)

Gewichtung

Wenn ausreichend klare Vermutungen oder empirische Hinweise bestehen, dass sich die Varianz des Fehlerterms im Modell umgekehrt proportional zu einer verfügbaren Variablen verhält, kann diese Variable zur Gewichtung der Schätzung verwendet werden und so deren Effizienz verbessern.³⁴ Für Modelle in absoluten Werten ist naheliegend, dass ein annähernd proportionaler Zusammenhang zwischen der Varianz und dem Quadrat jeder der in Frage kommenden Hilfsvariablen bestehen dürfte, weshalb beispielsweise der Kehrwert des Quadrats der Arbeitseinkommen als Gewicht in Betracht zu ziehen ist. Weiterhin lässt sich postulieren, dass sich durch den Einbezug des Hochrechnungsgewichts der WS in die Gewichtung der Schätzung allfällige Verzerrungen reduzieren lassen, die durch (a) die Überrepräsentation der grossen Einheiten in der WS-Stichprobe und (b) statistische Ausreisser (wie im vorhergehenden Unterabschnitt besprochen)³⁵ entstehen könnten. Deshalb werden – sowohl in den logarithmischen wie auch absoluten Modellen – auch Varianten unter Einbezug des WS-Hochrechnungsgewichts evaluiert.

Für die Evaluation werden folgende Varianten berücksichtigt:

- *Für die logarithmische Varianten:*
 - o **bm_ew_s23:** mit Gewicht = Hochrechnungsgewicht
 - o (Referenzvariante dazu, ohne Gewicht: **bm_s23**)

- *Für die Variante in Absolutwerten:*
 - o **bm_abs_ew:** mit Gewicht = Hochrechnungsgewicht / Hilfsvariable²
 - o (Referenzvariante dazu: **bm_abs**, mit Gewicht = 1 / Hilfsvariable²)

Vorgehen und Resultate

Um bei der Schätzung von logarithmischen Modellen Probleme zu vermeiden, wurden je nach Referenzjahr zwischen 10 und 14 Einheiten mit einem Wert von Null für den WS-Umsatz ausgeschlossen; in der Folge verbleiben im Schnitt pro Referenzjahr über 14 000 Einheiten mit verfügbaren Werten – was, insgesamt über die sieben Referenzjahre, eine Zahl von 100 223 Beobachtungen ergibt. Wie bereits erwähnt, liegt der Schwerpunkt der ersten Etappe im Vergleich der aus den verschiedenen Basisvarianten resultierenden Modellvorhersagen mit den direkten Schätzungen der WS. Somit liegt bei den Evaluationskriterien der Fokus auf QM-Aggregat. Letzterem liegt eine Zahl von 420 Beobachtungen (60 **noga2r**-Branchen in sieben Referenzjahren) zugrunde. Gezeigt werden jedoch, der Vollständigkeit halber, die Resultate sowohl für QM-Global als auch für QM-Aggregat. Ein Gesamtüberblick über alle evaluierten Basismodellvarianten und deren Resultate für die Qualitätsmasse ist in Tabelle T3 enthalten.

Logarithmische Modelle wurden eingehender analysiert als Modelle in absoluten Werten, und zwar aus den folgenden Gründen:

- Die Bestimmung einer optimalen Skalierung erwies sich als nicht trivial, weshalb nebst dem einfachsten (ohne Skalierung) logarithmischen Modell **bm** sechs Varianten, die sich von diesem ausschliesslich in Bezug auf die gewählte Skalierung unterscheiden, evaluiert wurden.
- Ein «Linear Mixed Model» (LMM) konnte ausschliesslich in der logarithmischen Formulierung erfolgreich implementiert werden, da LMM-Schätzungen in absoluten Werten nicht konvergierten, selbst nach Spezifikation diverser Vereinfachungen.³⁶
- Residuenplots (siehe Grafik G1) deuteten darauf hin, dass die Residuen in logarithmischen Modellen weniger stark von der idealen (normalen) Verteilung abweichen als in Modellen in absoluten Werten; insbesondere die ausgeprägt rechtsschiefe Verteilung der Residuen in letzteren dürfte ein Problem darstellen.

³³ Eine zielgerichtete Möglichkeit der Bereinigung um Ausreisser würde in der Anwendung von «robusten» Schätzverfahren bestehen, wie sie in diversen Statistikprogrammen zur Verfügung gestellt werden. Entsprechende Versuche mit der R-Package **robustlmm** (die robuste Schätzungen von LMM ermöglichen soll) waren jedoch wegen Konvergenzproblemen nicht erfolgreich. Es könnte dennoch lohnenswert sein, diesen Ansatz für weitere bevorstehende Modellierungen weiterzuverfolgen.

³⁴ Die im vorliegenden Dokument verwendete Formulierung von Gewichten w für lineare Modelle orientiert sich an der Spezifikation, wie sie von der

Statistiksoftware R verwendet wird; konkret, dass im Modell $\sum w\epsilon^2$ zu minimieren ist (und nicht $\sum (w\epsilon)^2$, wie in einigen anderen Programmen).

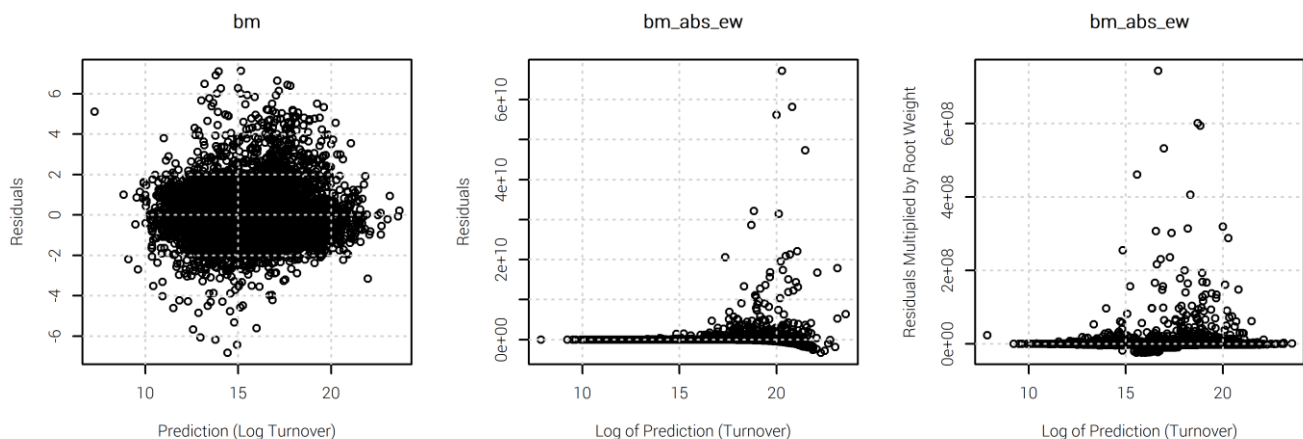
³⁵ Dies deshalb, weil die Konstruktion der Hochrechnungsgewichte der WS auch eine Robustifizierung beinhaltet mit dem Zweck, Ausreisser zu identifizieren und herunterzugewichten. Siehe «Methodenbericht Wertschöpfungsstatistik – Revision 2009: Statistische Datenaufbereitung und Hochrechnung», BFS, Neuchâtel, 2014

³⁶ LMM-Varianten werden erst in der zweiten Etappe (Detailmodelle, im nächsten Abschnitt) evaluiert und präsentiert.

Tabelle T3: Übersicht über die Basismodellvarianten und deren Evaluationsresultate

Modell-name	bm	bm_s0	bm_sl1	bm_sl23	bm_sl23a	bm_sl24	bm_sl24a	bm_sl23	bm_sl23	bm_sl23	bm_abs	bm_abs_sl23	bm_abs_sl23
Formulierung	logarithmisch										absolute Werte		
Umsatz / Beschäftigung	exponentiell							linear	exponentiell		linear		
Gewichtung	keine							Hochr.-Gewicht	keine		1 / VZÄ²	Hochr.-Gewicht / VZÄ²	
Ausreisserbereinigung	keine								Aus-schluss von Beob. mit Abso-lutwert des Residuums im gröss-ten Perzentil		keine		Aus-schluss von Beob. mit Abso-lutwert des standard. Residuums im gröss-ten Perzentil
Skalierung (nur für loga-rithmische Modelle)	keine	homo-gen, kein Aus-schluss von Extremwer-ten	heterogen (noga2r), kein Aus-schluss von Extremwer-ten	heterogen (noga2r), Ausschluss von Extremwerten. Als Extremwerte werden betrachtet: Beobachtungen mit Residuen mit...				... > 3 Standardabw.					
				... > 3 Standard-abw.	... > 3 Standard-abw. (nach noga2r)	... > 4 Standard-abw.	... > 4 Standard-abw. (nach noga2r)						
QM-Global (RMSE zwischen modellbasierten Schätzungen und erhobenen Werten, in CHF Mio, N = 100'223)													
Ungewichtet	1676.1	1674.5	1671.2	1668.8	1665.6	1667.6	1665.0	1682.3	1692.3	1673.5	1660.2	1659.2	1687.6
Gew. (HR)	115.7	117.3	116.6	115.8	115.6	115.9	115.8	116.2	123.2	116.0	116.9	115.5	116.9
Gew. (HR, Branchen)	64.3	67.3	68.2	65.1	65.4	65.8	66.3	65.9	73.2	65.4	69.6	65.6	66.3
QM-Aggregat (z-Werte der auf Ebene noga2r aggregierten Schätzungen, M = 420)													
Mittelwert	-1.129	2.681	0.947	0.356	0.323	0.576	0.628	0.760	0.758	0.196	2.827	0.879	0.563
RMS	1.521	4.472	1.675	1.030	1.091	1.213	1.392	1.535	1.261	0.930	5.346	2.082	1.782

Grafik G1: Residuenplots eines logarithmischen Modells (bm) und eines Modells in absoluten Werten (bm_abs_ew), Referenzjahr 2020



Erläuterungen: Für das Modell in absoluten Werten (**bm_abs_ew**) werden die Residuen sowohl in ihren Originalwerten (Mitte) als auch multipliziert mit der Quadratwurzel des für das Modell verwendeten Gewichts (rechts) dargestellt. Dies, da die bedingte Heteroskedastizität (im mittleren Plot deutlich erkennbar) eines der Motive für die gewählte Gewichtung ist. Diese bedingte Heteroskedastizität stellt im Hinblick auf die Effizienz theoretisch nur dann ein Problem dar, wenn sie nicht mittels einer geeigneten Gewichtung (das gewichtete Modell minimiert $\sum w_i^2$, also die Summe des Produkts aus Gewicht und Quadrat des Residuals) überwunden wird. Wie der Plot rechts zeigt, lässt sich das Heteroskedastizitätsproblem mit der gewählten Gewichtung weitgehend überwinden.

Als «Baseline»-Modell dient ein möglichst einfaches, logarithmisch formuliertes Modell namens **bm**. Möglichst einfach bedeutet insbesondere: keine erzwungene lineare Homogenität zwischen Beschäftigung und Umsatz (das heisst, ein exponentieller Zusammenhang ist zulässig), keine Gewichtung, keine Ausreisserbereinigung. Mittels dieses Baseline-Modells wurden dann diverse Varianten einer Skalierung für die Rücktransformation der Schätzwerte von logarithmierten zu absoluten Werten evaluiert, die beträchtliche Unterschiede offenbarten. Wird keinerlei Skalierung implementiert (**bm**), resultiert eine systematische Unterschätzung gegenüber den direkten Schätzungen (sichtbar im Mittelwert von QM-Aggregat: -1.129). Wird, wie im Modell **bm_s0**, eine möglichst einfache Skalierung, nämlich mit dem Faktor $\exp(s^2/2)$ gewählt (wobei s^2 das Quadrat des vom Modell gelieferten residualen Standardfehlers ist), resultiert hingegen eine Überschätzung (Mittelwert QM-Aggregat 2.681). Als wesentliche Probleme dieses Skalierungsansatzes werden vermutet:

- branchenbezogene bedingte Heteroskedastizität: Wird, anstelle eines einheitlichen Skalierungsfaktors, ein für die Branchen gemäss **noga2r** spezifischer Faktor verwendet (Modell **bm_s1**), verringert sich die Überschätzung deutlich; sie liegt jedoch nach wie vor deutlich über null (0.947).
- die deutliche Überkurtosis³⁷ der Schätzfehler: Eine weitere Verringerung der Überschätzung wurde erzielt mittels einer um Extremwerte bereinigten Berechnung der Skalierungsfaktoren (Modelle **bm_s23**, **bm_s23a**, **bm_s24** und **bm_s24a**). Diese Bereinigung wurde implementiert, indem jene Beobachtungen identifiziert und von der Berechnung der Skalierungsfaktoren ausgeschlossen wurden, deren Absolutwerte der Modellresiduen
 - o grösser als die geschätzte residuale Standardabweichung des Modells s mal drei (**bm_s23**) bzw. vier (**bm_s24**) waren;
 - o grösser als die innerhalb der jeweiligen **noga2r** geschätzte residuale Standardabweichung des Modells mal drei (**bm_s23a**) bzw. vier (**bm_s24a**) waren.

Dabei verringert sich die in **bm_s1** beobachtete Überschätzung (0.947) auf Werte zwischen 0.323 (**bm_s23a**) und 0.628 (**bm_s24a**).

Diese Unter- und Überschätzungen (reflektiert in den Mittelwerten von QM-Aggregat) weichen bei allen der bisher betrachteten Skalierungsvarianten deutlich von null ab. Bezüglich der Streuung von QM-Aggregat (gemessen mittels dem RMS) sorgen diese verfeinerten Varianten der Skalierung hingegen für eine deutliche

Verbesserung: In **bm_s23** fällt der RMS von QM-Aggregat auf 1.030 und weicht nicht mehr massiv von eins ab. Wird die nach wie vor bestehende Verzerrung für einen Moment ausser Betracht gelassen, lässt sich somit sagen, dass sich die Variabilität der Modellschätzungen gegenüber den direkten Schätzungen unter Verwendung dieser Skalierungsvarianten in einem akzeptablen Rahmen bewegt.

Bezüglich QM-Global (gewichtete Varianten) schneiden alle Modelle mit Skalierung schlechter ab als die Baseline-Variante. Da der Fokus in der ersten Etappe jedoch auf QM-Aggregat liegt, und da insbesondere für **bm_s23** und **bm_s23a** die Verschlechterung in QM-Global nicht massiv erscheint, ist die Skalierung dennoch als sinnvoll zu erachten. Für die noch zu betrachtenden Erweiterungen des Baseline-Modells wird die Problematik der systematischen Überschätzung (QM-Aggregat, Mittelwert) vorerst ignoriert³⁸ und der Skalierungsansatz **s23** beibehalten. Dieser schneidet zwar nur bedingt besser ab als **s23a**, jedoch wird **s23** als robuster betrachtet bezüglich zufälliger Häufungen von Extremwerten innerhalb einzelner **noga2r**-Branchen.³⁹ Die weiteren evaluierten Modelle lassen sich bezüglich ihrer Spezifikationen und Resultate wie folgt beschreiben:

- Lineare Homogenität Beschäftigung-Umsatz erzwungen (**bm_cs_23**): Daraus resultiert eine Verschlechterung (gegenüber **bm_s23**) aller Qualitätsmasse.
- Einbezug des Hochrechnungsgewichts der Wertschöpfungsstatistik in die Modellschätzung (**bm_ew_23**): Daraus resultiert eine deutliche Verschlechterung (gegenüber **bm_s23**) aller Qualitätsmasse.
- Bereinigung um Ausreisser, d.h. um Beobachtungen, deren Absolutwert des Residuums gemäss Initialmodell dem Perzentil mit den grössten Abweichungen angehört (**bm_oc_s23**): Während sich bei QM-Global kaum Änderungen ergeben gegenüber **bm_s23**, zeigen sich in QM-Aggregat Verbesserungen.

Modelle in absoluten Werten erzielten tendenziell schlechtere Resultate als logarithmische Modelle. Gegenüber der grundlegenden Variante des Modells in absoluten Werten (**bm_abs**)⁴⁰ verbesserte der Einbezug der Hochrechnungsgewichte (**bm_abs_ew**) die Prognoseeigenschaften zwar bezüglich der meisten Qualitätsmasse (im Falle von QM-Aggregat sehr deutlich), und auch die Implementierung einer Ausreisserbereinigung (**bm_abs_ew_oc**) scheint weitere geringfügige Verbesserungen zu bringen. Dennoch schneidet die favorisierte logarithmischen Modellvariante **bm_s23** tendenziell deutlich besser ab.

³⁷ Die Überkurtosis, hier gemessen als $k_{exc} = \frac{1}{n} \sum \varepsilon^4 / \left(\frac{1}{n} \sum \varepsilon^2 \right)^2 - 3$, gibt Aufschluss darüber, ob Werte mit extremen Abweichungen vom Mittelwert übermässig stark vertreten sind, bezogen auf deren in einer Normalverteilung zu erwartenden Häufigkeit. Für das Referenzjahr 2020 und das Modell **bm** beträgt $k_{exc} = 7.9$.

³⁸ Wie sich weiter unten (Evaluation der Detailmodelle) zeigen wird, relativiert sich die Problematik dieser Überschätzung nach Einbezug der weiteren Hilfsvariablen MWST-Umsatz markant.

³⁹ Dies, da der Schwellenwert (geschätzte residuale Standardabweichung des Modells mal drei) zum Ausschluss von Extremwerten bei $s \geq 3$ innerhalb der gesamten Population ermittelt wird und nicht innerhalb der jeweiligen **noga2r**.

⁴⁰ In dieser Variante des Modells in absoluten Werten (**bm_abs**) ist als Gewicht der Kehrwert der Hilfsvariablen (Beschäftigung in VZÄ) spezifiziert. Eine Modellvariante ohne Gewichtung wurde ebenfalls berechnet, jedoch wegen massiver Probleme bezüglich bedingter Heteroskedastizität nicht weiterverfolgt.

Fazit: Für die zweite Etappe sind logarithmischen Spezifikationen mit Skalierung (Varianten **bm_s23** oder **bm_s23a**) zu priorisieren. Zudem sind Modellvarianten mit Bereinigung um Ausreisser in Betracht zu ziehen, während die anderen der soeben evaluierten Erweiterungen wenig erfolgversprechend erscheinen.

3.4 Zweite Etappe (Detailvarianten)

Grundsätzliche Überlegungen

Die zweite Etappe (Evaluation der Detailmodellvarianten) beinhaltet den Einbezug zusätzlicher Datenquellen (ausserhalb der WS) und befasst sich mit den folgenden Modelleigenschaften, wobei als Ausgangspunkt die gemäss Evaluation der Basisvarianten favorisierte Grundstruktur von **bm_s23** und **bm_s23a** dient:

- Simultaner **Einbezug von zwei Hilfsvariablen**
- Einbezug von **binären erklärenden Variablen**
- Wahl der zweiten Hilfsvariablen: **Beschäftigung in VZÄ oder Arbeitseinkommen**
- Abschliessende Wahl der **Skalierungsvariante**: s23 oder s23a
- Abschliessende Wahl bezüglich **Bereinigung um Ausreisser**
- **«Fixed»- versus «Random»-Effekte** (konventionelle lineare Modelle versus «Linear Mixed Models») und, damit zusammenhängend, **Detaillierungsgrad der NOGA**
- Beurteilung des **Nutzens der «hybrid separat/gepoolten» Schätzung**
- Behandlung der **Beobachtungen mit erhobenem Wert von null** für die Zielvariable

In Analogie zur Präsentation der ersten Etappe bietet Tabelle T4 einen Überblick über die Detailmodellvarianten, die im Folgenden erläutert werden. Auf Grund der Anzahl der evaluierten Varianten und zwecks besserer Überschaubarkeit enthält Tabelle T4 nur eine Auswahl an Varianten; einen Überblick zu den Resultaten der Qualitätsmasse aller Detailmodellvarianten (ohne schematische Darstellung der Modellcharakteristika) liefert Tabelle TA1 im Anhang.

Einbezug von zwei Hilfsvariablen

Von besonderem Interesse für die Detailmodellvarianten ist die Hilfsvariable MWST-Umsatz. Wie im Abschnitt 1.2 dargelegt wurde, führen diverse Faktoren dazu, dass der MWST-Umsatz von unserer Zielvariablen «Umsatz gemäss WS» – trotz einer grundsätzlich stark ausgeprägten Korrelation – abweicht. Die Natur dieser Faktoren (hier sind insbesondere die von der MWST ausgenommenen Leistungen zu nennen) legt nahe, dass zwi-

schen der Ausprägung dieser Faktoren und der wirtschaftlichen Tätigkeit der Einheit ein enger Zusammenhang besteht. Somit ist die wirtschaftliche Tätigkeit einer Einheit, die im NOGA-Code erfasst wird, auf geeignete Art und Weise in die Modellierung einzu beziehen, insbesondere bezüglich der potentiell sehr unterschiedlichen Aussagekraft der MWST-Umsätze nach NOGA-Arten.

Der im Folgenden gewählte Ansatz begegnet diesem Problem, indem als Hilfsvariablen simultan (a) MWST-Umsätze und (b) Arbeitseinkommen gemäss STATENT (oder alternativ VZÄ gemäss STATENT, siehe Erläuterungen auf der nächsten Seite) berücksichtigt werden.⁴¹ Eine flexible Spezifikation ermöglicht, den jeweiligen Einfluss dieser beiden Hilfsvariablen auf die Vorhersage der Zielvariablen differenziert nach der wirtschaftlichen Tätigkeit gemäss empirischen Kriterien zu bestimmen. Die konkrete Umsetzung und die Überlegungen dazu werden im Technischen Anhang im Abschnitt «Grundlagen der Spezifikation für die Modellierung des WS-Umsatzes» dargelegt. Um auch Schätzwerte für jene Einheiten der Grundgesamtheit der Wertschöpfungsstatistik, für die keine MWST-Umsätze zur Verfügung stehen, berechnen zu können, wird als Lösungsansatz die «hybrid separat/gepoolte» Schätzung gewählt. Diese wird ebenfalls im Technischen Anhang detailliert beschrieben (im Abschnitt «Behandlung von Einheiten mit fehlendem MWST-Umsatz»). Für die Modellvorhersage wird somit auf zwei Teilmodelle zurückgegriffen: ein erstes für diejenige Einheiten, die über MWST-Umsätze verfügen, und ein zweites für Einheiten ohne MWST-Umsätze. Der empirische Nutzen dieses relativ komplexen Ansatzes gegenüber einfacheren Varianten, in denen auch im ersten Teilmodell ausschliesslich auf eine einzelne Hilfsvariable zurückgegriffen wird, lässt sich beziffern, indem solche vereinfachten Varianten ebenfalls geschätzt und in die Evaluation mit einbezogen werden.

Für die Evaluation werden folgende Varianten berücksichtigt:

- **dms_s23**: im Teilmodell mit MWST-Umsätzen: simultane Verwendung der MWST-Umsätze und der Arbeitseinkommen als Hilfsvariablen; im Teilmodell ohne MWST-Umsätze: Verwendung der Arbeitseinkommen als einzige Hilfsvariable
- **dms_0onlyvat_s23**: im Teilmodell mit MWST-Umsätzen: Verwendung der MWST-Umsätze als einzige Hilfsvariable; im Teilmodell ohne MWST-Umsätze: Verwendung der Arbeitseinkommen als einzige Hilfsvariable
- **dms_onlyinc_s23**: in beiden Teilmodellen: Verwendung der Arbeitseinkommen als einzige Hilfsvariable

Alle diese Varianten beinhalten:

- **binäre erklärende Variablen** für die Zugehörigkeit zu einer multinationalen Unternehmensgruppe und für Exporte (siehe nächster Unterabschnitt)

⁴¹ Die Anzahl der Beschäftigten wird hier somit als potentielle Hilfsvariable nicht in Betracht gezogen, da der effektive Arbeitsinput durch die VZÄ genauer erfasst

wird, und da letztere eine höhere Korrelation mit der Zielvariablen aufweisen (siehe Abschnitt 1.2).

- eine heterogene (nach **noga2r**) Skalierung, basierend auf $E(\hat{\epsilon}^2)/2$ als Skalierungsfaktor, nach Ausschluss der Beobachtungen mit einem Absolutwert von $\hat{\epsilon}$ grösser als die geschätzten residualen Standardabweichung des Modells mal drei

Binäre erklärende Variablen

Wie weiter oben erwähnt, beinhalten die Modellvarianten der zweiten Etappe einige binäre erklärende Variablen. Diese erfassen die folgenden Sachverhalte:

- die Zugehörigkeit der Einheit zu einer multinationalen Unternehmensgruppe (gemäss den Kriterien der Statistik der Unternehmensgruppen),
- das Vorhandensein von Exporten (gemäss Aussenhandelsstatistik des Bundesamtes für Zoll und Grenzsicherheit BAZG),
- das Vorhandensein von Exporten, deren Wert den Wert des Umsatzes gemäss MWST übersteigt

Der Nutzen dieser binären erklärenden Variablen für die Präzision der Modellvorhersage lässt sich beziffern, indem eine Modellvariante ohne diese Variablen geschätzt wird.

Für die Evaluation wird folgende Variante berücksichtigt:

- **dms_p_s23**: ohne binäre erklärende Variablen für Zugehörigkeit zu einer multinationalen Unternehmensgruppen und für Exporte
- (Referenzvariante dazu, mit Einbezug der binären erklärenden Variablen: **dms_s23**)

Hilfsvariablen Beschäftigung oder Arbeitseinkommen

Die STATENT bzw. die AHV-Daten stellen zwei potentiell relevante Hilfsvariablen zur Verfügung: Die Beschäftigung in Vollzeitäquivalenten (VZÄ) sowie die Arbeitseinkommen. Erstere ist konzeptionell vergleichbar mit den in der ersten Etappe verwendeten VZÄ des Rahmens der WS und somit bereits erprobt. Demgegenüber weist die Variable Arbeitseinkommen gewisse theoretische Vorzüge auf:

- Sie erfasst – wie die Zielvariable Umsatz – die Summe einer während eines Kalenderjahres erbrachten Leistung, während sich die VZÄ der STATENT jeweils nur auf den Monat Dezember eines Jahres beziehen.
- Sie misst wie die Zielvariable WS-Umsatz eine monetäre Leistung.

- Sie wird direkt erhoben und nicht mittels eines statistischen Modells berechnet, wie dies bei den VZÄ der STATENT grösstenteils der Fall ist. Bei der STATENT sind die VZÄ nur für die Unternehmen, die an Erhebungen (BESTA, Profiling oder Profiling Light) teilgenommen haben – mehrheitlich Mehrbetriebsunternehmen – direkt bekannt, also Originalwerte. Für die anderen Unternehmen werden die VZÄ eingesetzt, gestützt auf einem Modell, das für die Einzelbetriebsunternehmen mit aus einer Erhebung (Beschäftigungsstatistik BESTA, oder Aktualisierungserhebung des Betriebs- und Unternehmensregisters ERST) bekannten VZÄ und vorhandenen Arbeitseinkommen gemäss AHV geschätzt wurde.

Um eine empirisch fundierte Beurteilung zu ermöglichen, welche der beiden Variablen bessere Vorhersageeigenschaften aufweist, werden folglich Modellvarianten sowohl mit VZÄ gemäss STATENT als auch mit Arbeitseinkommen als Hilfsvariablen (nebst den MWST-Umsätzen) geschätzt.⁴²

Für die Evaluation wird folgende Variante berücksichtigt:

- **dms_fte_s23**: mit MWST-Umsatz und VZÄ der STATENT als Hilfsvariablen
- (Referenzvariante dazu, mit MWST-Umsatz und Arbeitseinkommen als Hilfsvariablen: **dms_s23**)

Skalierungsvariante

In der ersten Etappe wurde noch keine abschliessende Wahl zwischen den beiden Skalierungsvarianten s23 und s23a getroffen, da diese sehr ähnliche Resultate erzielten. Die deutlich verbesserte Präzision der Schätzungen in der zweiten Etappe (sichtbar im Vergleich der Resultate für QM-Global zwischen Tabellen T3 und T4), insbesondere als Folge des Einbezugs der MWST-Umsätze als Hilfsvariable, könnte auch auf die Eignung dieser beiden Varianten einen Einfluss ausüben. Deshalb werden im Rahmen der zweiten Etappe beide Varianten nochmals evaluiert.

Für die Evaluation werden folgende Varianten berücksichtigt:

- **dms_s23a**: heterogene (nach **noga2r**) Skalierung, Verwendung von $E(\hat{\epsilon}^2)/2$ als Skalierungsfaktor, nach Ausschluss der Beobachtungen mit einem Absolutwert von $\hat{\epsilon}$ grösser als die innerhalb der jeweiligen **noga2r** geschätzten residualen Standardabweichung des Modells mal drei
- (Referenzvariante dazu, ebenfalls heterogene Skalierung nach **noga2r** unter Verwendung von $E(\hat{\epsilon}^2)/2$ nach Ausschluss der Beobachtungen mit einem Absolutwert von $\hat{\epsilon}$ grösser als die geschätzten residualen

⁴² Auf die Schätzung eines Modells mit simultanem Einbezug der Hilfsvariablen VZÄ und Arbeitseinkommen (nebst dem MWST-Umsatz) wurde aufgrund der

ausgeprägten Kollinearität zwischen diesen beiden Variablen verzichtet (Korrelationskoeffizient der Logarithmen der beiden Variablen: 0.971).

*Standardabweichung des Modells mal drei, dies jedoch ohne Differenzierung nach **noga2r**: **dms_s23**)*

Bereinigung um Ausreisser

Da die erste Etappe keine eindeutigen Hinweise für oder wider eine Ausreisserbereinigung lieferte, werden in der zweiten Etappe nochmals Varianten mit und ohne Ausreisserbereinigung gerechnet und miteinander verglichen.

Für die Evaluation werden die folgenden Varianten berücksichtigt:

- **dms_oc_s23**: wiederholte Schätzung des Modells **dms_s23**, nach Ausschluss jener Beobachtungen, deren Absolutwert des Residuums dem Perzentil mit den grössten Schätzfehlern angehört
- (Referenzvariante dazu, ohne Ausreisserbereinigung: **dms_s23**)

«Fixed»- versus «Random»-Effekte und Detaillierungsgrad der NOGA

Wie bereits mehrfach erwähnt, soll die Modellierung Heterogenität der Parameter in Bezug auf Gruppierungen von Einheiten – konkret gemäss derer wirtschaftlichen Tätigkeit – zulassen. Dies lässt sich mittels Gruppeneffekten spezifizieren, und zwar entweder als «Fix»-Effekte (und somit im Rahmen eines konventionellen OLS-Ansatzes) oder als «Random»-Effekte. So genannte «Linear Mixed-Modelle» (LMM) erlauben, im gleichen Modell gewisse Effekte als Fix und andere als Random zu spezifizieren, wobei Random-Effekte auch verschachtelt («nested») werden können.⁴³

Aus theoretischer Sicht sollten folgende Bedingungen erfüllt sein, damit für eine Gliederungsebene die Spezifikation als Random-Effekt sinnvoll ist (in dem Sinne, dass Effizienzgewinne zu erwarten sind, ohne dass das Risiko von Verzerrungen zu gross wird):

- eine gewisse **Mindestanzahl von Ausprägungen**, da die Schätzung des Modells jeweils eine Schätzung der Varianz des Random-Effekts beinhaltet. Letzteres ist bei einer geringen Anzahl von Ausprägungen problematisch, wobei es hier keine klar definierten Schwellenwerte oder Kriterien gibt.
- das Vorhandensein (bzw. das Risiko des Vorhandenseins) von **Ausprägungen mit einer geringen Anzahl von Beobachtungen**. Je grösser die Anzahl der Beobachtungen einer bestimmten Ausprägung, desto mehr nähern sich die Resultate der Random- jenen der Fix-

Spezifikation an; das heisst, das Argument zugunsten der Random-Spezifikation – nämlich, dass das Modell zufallsbedingte Ausschläge in den strichprobenseitig beobachteten Effekten tendenziell abmildert – wird irrelevant. Auch hier gibt es keine klar definierten Schwellenwerte.

Im Kontext dieser Überlegungen kann als Entscheidungsgrundlage die folgende Tabelle dienen. Sie zeigt, für drei in Frage kommende Stufen von Branchengliederungen (**noga2r**, **noga3** und **noga6**)⁴⁴, jeweils die Anzahl derer Ausprägungen, die Anzahl der Beobachtungen in den Ausprägungen mit der kleinsten und mit der grössten Anzahl von Beobachtungen sowie die Anzahl Ausprägungen mit lediglich einer vorhandenen Beobachtung (da die Resultate zwischen Referenzjahr schwanken, sind wo relevant Spannen von Werten ausgewiesen):

Branchengliederung	Anzahl der Ausprägungen:		Anzahl der Beobachtungen:	
	total	mit nur einer Beobachtung	minimal	maximal
noga2r	70	0	27–31	1575–1663
noga3	224–226	5–8	1	464–534
noga6	622–636	46–60	1	264–534

Auf der Basis all dieser Überlegungen wurden die in der folgenden Tabelle erscheinenden Varianten von Spezifikationen in Betracht gezogen. Für die Branchengliederungen **noga3** und **noga6** sind eher Random-Effekte angezeigt, da Ausprägungen mit sehr wenigen Beobachtungen (im Extremfall mit nur einer einzigen) durchwegs vorhanden sind.⁴⁵ Weniger klar erscheint der Fall bei **noga2r**, wo Fix- und Random-Effekte a priori ähnlich sinnvoll sein können. Als Referenzvariante dient ein Modell mit ausschliesslich Fix-Effekten. Bei den weiteren Varianten handelt es sich um «Mixed»-Spezifikationen, wobei die Variante r1 «genestete» Random-Effekte (**noga2r**, **noga6**) beinhaltet.

Variante	Fix-Effekt(e)	Random-Effekt(e)
(Referenzvariante)	noga2r	
r0	(keine)	noga2r
r1	(keine)	noga2r, noga6
r2	noga2r	noga3
r3	noga2r	noga6

Da durch die Berücksichtigung von «Random»-Effekten detaillierter als **noga2r** in gewissen Varianten Konvergenzprobleme auftraten, wurde für die Spezifikation dieser Varianten auf eine linear homogene Formulierung zwischen abhängiger Variable und Hilfsvariablen ausgewichen. Die Spezifikation vereinfacht sich

⁴³ Zur Schätzung von LMM wird auf das R-Package **lme4** zurückgegriffen.

⁴⁴ Bei **noga3** und **noga6** handelt es sich um die ersten drei Stellen (Gruppe) bzw. alle sechs Stellen (Art) der NOGA-Klassifizierung. Siehe <https://www.bfs.admin.ch/bfs/de/home/statistiken/industrie-dienstleistungen/nomenklaturen/noga.html> für mehr Informationen.

⁴⁵ Die Verwendung von Branchengliederungen detaillierter als **noga2r** ist potentiell sinnvoll, da innerhalb vieler NOGA-Abteilungen mit einer grossen Heterogenität zu rechnen ist. Ein Beispiel dafür ist die NOGA-Abteilung 68 (Immobilienwesen), in welcher gewisse Arten hauptsächlich Leistungen anbieten, die von der MWST ausgenommen sind (Vermietung), andere Arten jedoch mehrheitlich MWST-pflichtige Umsätze erwirtschaften (Vermittlung und Verwaltung).

dadurch (die Anzahl der zu schätzenden Random-Effekte halbiert sich), und Konvergenzprobleme konnten so vermieden werden. Bei der im Kontext der Basisvarianten in der ersten Etappe evaluierten linear homogenen Formulierung zeigte sich zwar eine Verschlechterung der Modelleigenschaften; jedoch ist denkbar, dass für die Detailmodelle dank des umfangreicheren Informationsgehalts der verwendeten Hilfsvariablen die Unterstellung eines linear homogenen Zusammenhangs weniger gravierende Auswirkungen hat.

Für die Evaluation werden folgende Varianten berücksichtigt:

- **dms_r0_s23:** *noga2r* als «Random»-Effekte, für beide Teilmodelle
- **dms_r10_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r/noga6* als genestete «Random»-Effekte
 - o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Random»-Effekte
- **dms_r10_0cs_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r/noga6* als genestete «Random»-Effekte und lineare Homogenität zwischen abhängiger Variable und Hilfsvariablen
 - o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Random»-Effekte
- **dms_r20_0cs_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r* als Fix-Effekte, *noga3* als «Random»-Effekte und lineare Homogenität zwischen abhängiger Variable und Hilfsvariablen
 - o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Fix»-Effekte
- **dms_r30_0cs_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r* als Fix-Effekte, *noga6* als «Random»-Effekte und lineare Homogenität zwischen abhängiger Variable und Hilfsvariablen
 - o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Fix»-Effekte
- **dms_r10_0cap_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r/noga6* als genestete «Random»-Effekte; zudem Bereinigung unplausibler Schätzwerte der Koeffizienten (siehe Techn. Anhang)
 - o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Random»-Effekte
- **dms_r10_0cs_0cap_s23:**
 - o für das Teilmodell mit MWST-Umsätzen: *noga2r/noga6* als genestete «Random»-Effekte und lineare Homogenität zwischen

abhängiger Variable und Hilfsvariablen; zudem Bereinigung unplausibler Schätzwerte der Koeffizienten (siehe Techn. Anhang)

- o für das Teilmodell ohne MWST-Umsätze: *noga2r* als «Random»-Effekte
- (Referenzvariante dazu, mit *noga2r* als Fix-Effekte, für beide Teilmodelle: **dms_s23**)

Nutzen der «hybrid separat/gepoolten» Schätzung

Der im Technischen Anhang formell beschriebene «hybrid separat/gepoolte» Ansatz ist dadurch motiviert, dass damit für die Subpopulation der Unternehmen ohne verfügbare MWST-Umsätze Effizienzgewinne möglich sind gegenüber einem separaten Ansatz. Ob sich solche Effizienzgewinne tatsächlich manifestieren, lässt sich ansatzweise verifizieren, indem für eine der bis hierhin präsentierten Modellvarianten (die sich alle auf den «hybrid separat/gepoolten» Ansatz stützen) eine vereinfachte Alternativvariante gerechnet wird, in der das Teilmodell für die Unternehmen ohne MWST-Umsätze isoliert (also separat), jedoch mit ansonsten unveränderter Spezifikation geschätzt wird.

Für die Evaluation werden die folgenden Varianten berücksichtigt:

- **dms_r10_0cap_1x_s23:** Teilmodell ohne MWST-Umsätze isoliert geschätzt
- (Referenzvariante dazu, Teilmodell ohne MWST-Umsätze «gepoolt» gemäss Gleichung (6) im Technischen Anhang: **dms_r10_0cap_s23**)

Beobachtungen mit erhobenem Wert von null für Zielvariable

Um in der Schätzung den Ausschluss von Beobachtungen mit einem erhobenem Wert von Null für den Umsatz gemäss WS zu vermeiden (siehe Abschnitt 3.3; je nach Referenzjahr 10 bis 14 betroffene Beobachtungen), wurde auch eine Modellvariante evaluiert, in der die abhängige Variable vor der Logarithmierung um einen geringfügigen konstanten Betrag erhöht wird (zur Gewährleistung der Konsistenz wird für die Berechnung der Vorhersagen dieser Betrag dann wieder subtrahiert).

Für die Evaluation werden die folgenden Varianten berücksichtigt:

- **dms_Logplus_r10_s23:** Spezifikation von $\log(U_i + 1000)$ für die abhängige Variable (Werte in Franken) zwecks Schätzung des Modells⁴⁶
- (Referenzvariante dazu: **dms_r10_s23**, mit $\log U_i$ als abhängige Variable)

⁴⁶ Der gewählte Erhöhungsbetrag von 1000 Franken lässt sich theoretisch dadurch rechtfertigen, dass in der Wertschöpfungsstatistik alle Kennzahlen auf 1000

gerundet erhoben werden. Zudem führte die Verwendung kleinerer Beträge (500, 100 oder 1) in beiden Teilmodellen jeweils zu deutlich höheren residualen Standardfehlern.

Vorgehen und Resultate

Tabelle T4 liefert den Gesamtüberblick über eine Auswahl der Detailmodellvarianten und deren Resultate für die Qualitätsmasse, während sich die Resultate aller evaluierten Detailmodellvarianten (ohne Erläuterung der Modellcharakteristika) im Anhang in Tabelle TA1 finden lassen. Ebenfalls im Anhang stellt Tabelle TA2 Resultate aufgeschlüsselt nach Referenzjahren und Tabelle TA3 aufgeschlüsselt nach NOGA-Abschnitt dar, diese jedoch

ausschliesslich für QM-Global gewichtet mittels Hochrechnungsgewicht und Branchenfaktoren.

Als Evaluationskriterium für die zweite Etappe steht QM-Global im Vordergrund. Dennoch ist zunächst ein Blick auf die Resultate für QM-Aggregat interessant: Gegenüber der ersten Etappe zeigen sich deutliche Verbesserungen in den mittleren z-Werten; diese weichen nun in den meisten der in Tabelle T4 aufgeführten Modellvarianten um weniger als 0.1 von null ab. Somit verschwindet

Tabelle T4: Übersicht über einige Detailmodellvarianten und deren Evaluationsresultate

Modell- name	dms_s23	dms_s23a	dms_fie_s23	dms_oc_s23	dms_r0_s23	dms_r10_s23	dms_r10_0cap_s23	dms_r10_46cs_0cap_s23	dms_r10_46cs_0cap_i_s23
Skalierung	Teilmodell mit MWST: heterogen (noga2r); Teilmodell ohne MWST: homogen. Ausschluss von Extremwerten. Als Extremwerte werden betrachtet: Beobachtungen mit Residuen...								
	... > 3 Standardabw.	... > 3 Standardabw. (nach noga2r)	... > 3 Standardabw.						
Verwendete Hilfsvariablen	Arbeitseinkommen, MWST-Umsätze		Vollzeitäquivalente, MWST-Umsätze	Arbeitseinkommen, MWST-Umsätze					
Ausreisser- bereinigung	keine			Ausschluss von Beob. mit Absolutwert des Residuums im gröss- ten Perzentil	keine				
Fix-Effekte	noga2r				keine				
Random-Effekte	keine				noga2r	Teilmodell mit MWST: noga2r/noga, Teilmodell ohne MWST: noga2r			
Bereinigung unplau- sibler Koeffizienten- Schätzwerte	keine						ja (siehe Technischer Anhang)		
Umsatz / Hilfsvari- ablen	Beide Teilmodelle: exponentiell						Teilmodell mit MWST: exponentiell, ausser NOGA 46: linear homogen; Teilmodell ohne MWST: exponentiell.		
Punktueller Korrekturen	keine								Ja (siehe Text)
QM-Global (RMSE zwischen modellbasierten Schätzungen und erhobenen Werten, in CHF Mio., N = 99'187) ¹									
Ungewichtet	615.2	650.3	614.1	631.6	614.3	600.1	600.6	569.6	569.6
Gew. (HR)	49.3	51.9	49.5	52.9	49.4	46.6	46.7	44.9	44.9
Gew. (HR, Bra.)	35.5	36.2	36.1	37.1	35.0	31.2	31.2	30.8	30.8
QM-Aggregat (z-Werte der auf Ebene noga2r aggregierten Schätzungen, M = 419) ²									
Mittelwert	-0.130	-0.017	-0.116	0.060	-0.067	-0.003	-0.018	-0.017	-0.106
RMS	1.829	1.955	1.945	2.089	1.892	1.902	1.870	1.872	1.624

¹ Die Anzahl der für QM-Global herangezogenen Beobachtungen (N = 99'187) ist geringer als bei der Evaluation der Basismodellvarianten (in Tabelle T3), da nur jene Beobachtungen der originalen WS-Nettostichprobe berücksichtigt wurden, die auch gemäss den aktuellsten Daten die Kriterien zur Zugehörigkeit zum WS-Rahmen erfüllen.

² Die Anzahl der für QM-Aggregat herangezogenen Beobachtungen (M = 419) ist um eins geringer als bei der Evaluation der Basismodellvarianten (in Tabelle T3), da für die Evaluation der Detailmodellvarianten eine Beobachtung ausgeschlossen wurde (NOGA 71 im Referenzjahr 2018), in der für eine Einheit bei den Hilfsvariablen unplausible Werte erhoben worden waren, was zu einem unerwünscht starken Einfluss auf QM-Aggregat geführt hätte. Für die betroffene Einheit wird im Rahmen der für die Statistikproduktion schlussendlich ausgewählten Modellvariante **dms_r10_46cs_0cap_i_s23** auf die Modellvorhersage eines Alternativmodells zurückgegriffen (siehe Erläuterungen auf der nächsten Seite).

das zuvor festgestellte Problem von systematisch höheren Schätzergebnissen in Bezug auf die direkten Schätzungen der WS weitgehend, als Folge des Einbezugs der MWST-Umsätze als Hilfsvariable. Hingegen ist die Streuung dieser z-Werte grösser geworden: Während die in der ersten Etappe favorisierten Modellvarianten RMS-Werte zeigten, die sich dem Wert von 1 annäherten oder diesen gar unterschritten, liegen die RMS für die Detailmodelle durchwegs über 1.5. Dies zeigt, dass – wie in Abschnitt 3.1 erläutert – durch den Einbezug von Daten ausserhalb des WS-Rahmens sich die Übereinstimmung zwischen Modellvorhersage und direkten Schätzungen verringert, und zwar weil diese Daten tatsächlich zu präziseren Schätzungen der Umsätze führen (siehe Erläuterungen im Anhang).

Der Ausschluss einer Reihe von potentiellen Alternativmodellen führt zu **dms_r10_0cap_s23** als grundsätzlich beste Variante. Detailliert erläutert und begründet wird das Vorgehen, das zu diesem Schluss geführt hat, im Anhang des vorliegenden Berichts. Um von dieser Variante zur schlussendlich für die Statistikproduktion verwendeten Variante zu gelangen, wurden noch zwei Aspekte identifiziert und implementiert, mit denen gezielte Verbesserungen möglich waren:

- Eine kurze Betrachtung nach NOGA-Abschnitten der Resultate für QM-Global (ausschliesslich in der mittels Hochrechnungsgewicht und Branchenfaktoren gewichteten Variante, Tabelle TA3) aller Modellvarianten brachte eine interessante Erkenntnis: Wird **dms_r10_0cap_s23** um das Erfordernis von linearer Homogenität erweitert (Variante **dms_r10_0cs_0cap_s23**), so schneiden einige NOGA-Abschnitte besser ab. Besonders umfangreich ist die Verbesserung im NOGA-Abschnitt «Handel» (G). Eine weiterführende Analyse zeigte, dass im Grosshandel (NOGA-Abteilung 46; Teil des Abschnitts G) das Erzwingen von linearer Homogenität in allen Referenzjahren für eine Verbesserung sorgte, während dies in den anderen Abteilungen des Abschnitts G nicht der Fall war. In der Folge wurde eine verfeinerte Variante spezifiziert und evaluiert, in welcher **lineare Homogenität ausschliesslich für die NOGA-Abteilung 46** erzwungen wird (Variante **dms_r10_46cs_0cap_s23**).

- **Punktueller Korrekturen** mittels Zurückgreifen auf Schätzwerte eines anderen Modells für einige wenige, gezielt identifizierte Beobachtungen, deren Schätzwerte von **dms_r10_46cs_0cap_s23** unplausibel hoch waren (siehe nächster Unterabschnitt). Das Ergebnis dieser Korrekturen wird als «Modell» **dms_r10_46cs_0cap_i_s23** bezeichnet.

Somit und zur Gewährleistung der Modellstabilität im Zeitverlauf wurde entschieden, für die **Statistikproduktion aller Referenzjahre** auf die **Modellvariante dms_r10_46cs_0cap_i_s23** zurückzugreifen.

Grafische Darstellung und punktueller Korrekturen

Die dritte der in Abschnitt 2.1 formulierten Anforderungen (Vermeidung unplausibler modellbedingter Ausschläge im Zeitverlauf) wurde gewährleistet, indem die im vorhergehenden Unterabschnitt favorisierte Modellvariante **dms_r10_46cs_0cap_s23** vertieft betrachtet und einzelne Beobachtungen identifiziert wurden, für die sich punktueller Anpassungen an den Modellresultaten aufdrängten. Dazu wurde die folgende zweigleisige Vorgehensweise gewählt:

- Aggregation der geschätzten Werte des Umsatzes nach NOGA-BFS50⁴⁷ (nach Einsetzung der erhobenen Werte für die Einheiten der Nettostichprobe, also analog dem Vorgehen für QM-Aggregat), und **Visualisierung der resultierenden Reihen** – zusammen mit den parallel dazu hochgerechneten Resultaten gemäss WS⁴⁸ und einer Auswahl von Alternativmodellen – im zeitlichen Verlauf (siehe Grafik G2)
- vertiefte Analysen einzelner Einheiten**, die nicht in der Nettostichprobe der WS enthalten sind und deren Modellvorhersage einen hohen Einfluss auf den Totalwert ihrer **noga2r**-Branche ausübt. Konkret wurden jene Einheiten analysiert, für die in einem oder mehreren Referenzjahren der vorhergesagte Umsatz 10% oder mehr des totalen Umsatzes (hochgerechnet gemäss WS) innerhalb der **noga2r**-Branche ausmachte.

⁴⁷ Die Reihen werden hier nach der Gliederungsstufe NOGA-BFS50 (die für die Veröffentlichung der Resultate im Rahmen der STAGRE massgebend ist) ausgewiesen und nicht nach **noga2r**. Dies deshalb, weil der höhere Detaillierungsgrad von **noga2r** zu einer Vielzahl von Datenpunkten führen würde, die aus Gründen der Datenvertraulichkeit nicht veröffentlicht werden dürften. Für interne Analysen wurden auch Reihen nach **noga2r** visualisiert und ausgewertet.

⁴⁸ Diese hochgerechneten Resultate gemäss WS entsprechen nicht exakt den direkten Schätzungen der WS, da im vorliegenden Kontext für die Branchenzuordnung der Einheiten aus Kohärenzgruppen die NOGA-Codes gemäss aktuellsten Quellen (und nicht gemäss Rahmen der WS) verwendet werden.

Diese Vorgehensweise erlaubte zudem auch, allfällige Schwächen des Modells (oder der Modelle) in Bezug auf bestimmte Gegebenheiten (z.B. Branchen; Art der Erhebung des MWST-Umsatzes) zu erkennen. Bezüglich der favorisierten Modellvariante **dms_r10_46cs_0cap_s23** zeigte sich, dass die Resultate im Allgemeinen zufriedenstellend waren. Die Analyse in (b) ermöglichte jedoch, unter den zwölf die Filterkriterien erfüllenden Einheiten deren sechs zu identifizieren, für die in einem oder mehreren Referenzjahren die Resultate als unplausibel zu erachten waren und stattdessen auf die Schätzwerte anderer Modelle zurückgegriffen wurde:

- für **fünf Einheiten** wurden für eines oder mehrere Referenzjahre die Schätzwerte des entsprechenden **Teilmodells ohne Berücksichtigung der MWST-Umsätze** eingesetzt;
- für **eine Einheit** wurde für ein Referenzjahr der Schätzwert des entsprechenden **Teilmodells mit ausschliesslicher Berücksichtigung des MWST-Umsatzes** eingesetzt.

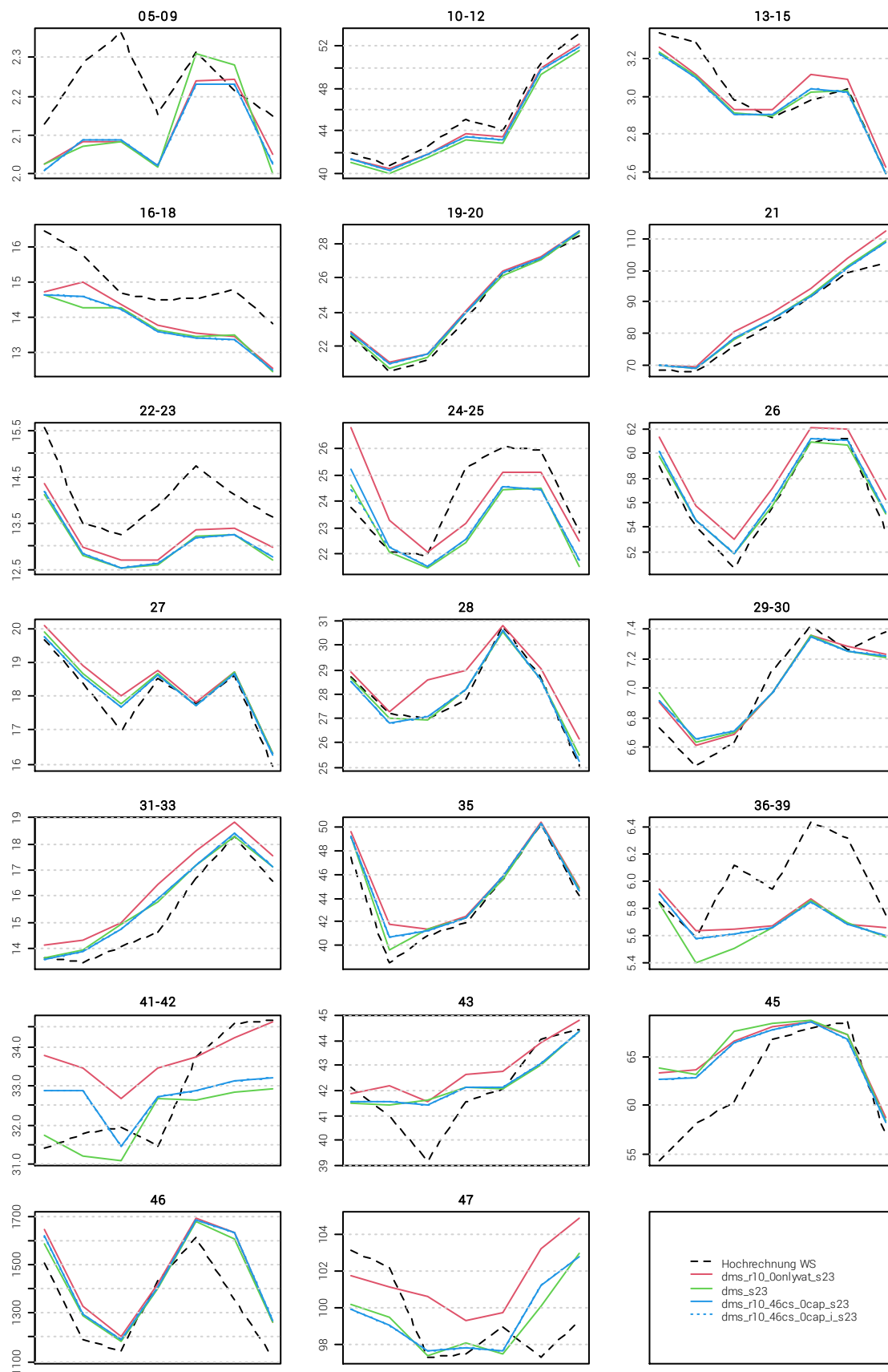
Die mutmasslichen Ursachen für die angezweifelte Eignung des favorisierten Modells in diesen Fällen waren vielfältig.⁴⁹ In der Summe lieferten diese Einzelfälle keine Hinweise auf systematische Schwächen des Modells. Die aus diesen punktuellen Einsetzungen resultierenden Vorhersagen werden als «Modell» **dms_r10_46cs_0cap_i_s23** bezeichnet; es handelt sich hier um die schlussendlich für die Statistikproduktion verwendeten Daten. Die daraus abgeleiteten aggregierten Resultate sind – nebst jenen von **dms_r10_46cs_0cap_s23** – in Grafik G2 dargestellt. Zur Illustration enthält Grafik G2 zudem die Resultate von **dms_r10_onlyvat_s23** (ausschliessliche Verwendung der MWST-Umsätze als Hilfsvariable, wo vorhanden) sowie von **dms_s23** (Verzicht auf «Mixed»-Spezifikation, also ausschliesslich Fix-Effekte).

In den Branchen des sekundären Sektors (NOGA-Abteilungen 05 bis 43) sticht keine der dargestellten Modellvarianten besonders hervor; weder bezüglich ihrer Abweichungen von den hochgerechneten Resultaten, noch bezüglich allfälliger unplausibler Ausschläge. Anders sieht es im tertiären Sektor (NOGA-Abteilungen 45 bis 96) aus: **dms_r10_onlyvat_s23** zeigt hier die Tendenz, stärker als die anderen Varianten von den hochgerechneten Resultaten abzuweichen, und neigt zudem etwas mehr zu erratischen Ausschlägen. Dies unterstreicht den Nutzen der Verwendung der Arbeitseinkommen als Hilfsvariable auch für jene Beobachtungen, für die MWST-Umsätze verfügbar sind. Rein visuell betrachtet erscheint die favorisierte Mixed-Spezifikation (**dms_r10_46cs_0cap_s23**) im Aggregat nicht näher an den hochgerechneten Resultaten als die Fix-Effekt-Spezifikation **dms_s23**; die Bevorzugung von Ersterer lässt sich allerdings wie weiter oben erläutert durch ihre besseren Resultate in QM-Global

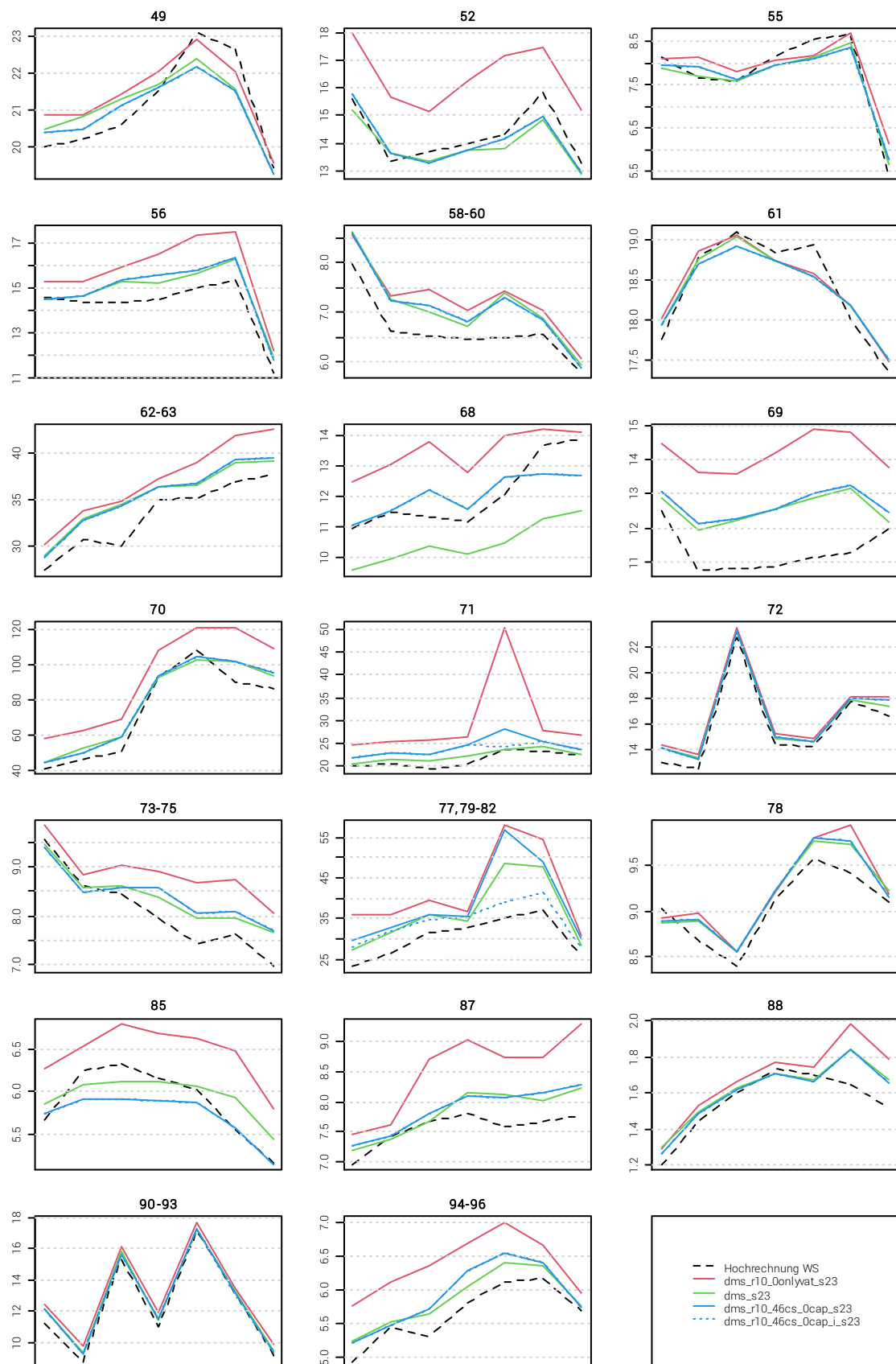
rechtfertigen. In **dms_r10_46cs_0cap_i_s23** zeigt sich schliesslich, dass die dort vorgenommenen punktuellen Anpassungen in einigen Branchen (71 sowie 77, 79–82) zu teils massiv plausiblen Resultaten führen.

⁴⁹ z.B. Erhebungsfehler beim MWST-Umsatz, zu hohe verteilte MWST-Umsätze in Gruppen mit untypischer Branchenzusammensetzung, problematische Änderungen der NOGA-Einteilung der Einheit über den Zeitverlauf

Grafik G2: Verlauf (2014 – 2020) der Umsätze nach NOGA-BFS-50 gemäss hochgerechneten Resultaten der WS sowie einigen Modellvarianten, in Mia. Franken (Fortsetzung auf der nächsten Seite)



Grafik G2 (Fortsetzung): Verlauf (2014 – 2020) der Umsätze nach NOGA-BFS-50 gemäss hochgerechneten Resultaten der WS sowie einigen Modellvarianten, in Mia. Franken



4 Abschliessende Bemerkungen und Ausblick

4.1 Abschliessende Bemerkungen

Der Modellierungsansatz für den Umsatz, wie er im vorliegenden Bericht begründet und beschrieben wird, hat massgeblich dazu beigetragen, den Informationsumfang und somit auch die Analysemöglichkeiten der Statistik der Unternehmensgruppen (STAGRE) zu erweitern. Grundlage dazu bildete die Verknüpfung von Daten aus statistischen Quellen (Wertschöpfungsstatistik) einerseits mit administrativen Quellen (Mehrwertsteuer, AHV-Register) andererseits. Hierzu konnte auf die systematischen Vorarbeiten aufgebaut werden, die im Bundesamt für Statistik und in Zusammenarbeit mit anderen Stellen wie der Eidgenössischen Steuerverwaltung und den AHV-Ausgleichskassen getätigt wurden, um diese Daten für statistische Zwecke nutzbar zu machen. Somit entspricht dieser Ansatz exakt dem folgenden, als Teil des zweiten strategischen Ziels des Statistischen Mehrjahresprogrammes 2020–2023 des Bundes formulierten Anspruch: «Mit bereits geschaffenen neuen Möglichkeiten wie beispielsweise der in den letzten Jahren systematisch weiterentwickelten Verknüpfung von Daten oder auch dem innovativen Einsatz von Methoden soll das Angebot an statistischen Informationen der Bundesstatistik weiter ausgebaut werden.» Erwähnenswert ist auch, dass für die hier erfolgte Erweiterung des Informationsangebots keine zusätzliche Belastung der Unternehmen (in der Form neuer oder ausgeweiteter Befragungen) erforderlich war.

Um für die Publikation von Resultaten des Umsatzes den grundlegenden Anforderungen des Datenschutzes gerecht zu werden, werden die üblichen Kriterien eingehalten: Es werden in den Tabellen keine Werte ausgewiesen für Zellen, in denen weniger als vier Einheiten positive Umsätze aufweisen, oder in denen die Einheit mit dem grössten Umsatz 60% oder mehr zum Zellentotal beiträgt. Trotz dieser Einschränkungen stehen in der STAGRE nun Resultate für den Umsatz zur Verfügung nach einer Vielzahl von Gliederungskriterien (Art der Unternehmensgruppe, Branchen, Rechtsform, Grössenklasse der Gruppe, Regionalisierungsgrad der Gruppe, Sitzland), die sich teilweise auch kombinieren lassen.⁵⁰

4.2 Ausblick

Die während der Erarbeitung des Umsatzes der STAGRE gewonnenen Erfahrungen und Erkenntnisse bilden eine umfangreiche und flexible Grundlage, auf der weitere Entwicklungen aufbauen können. Als nächster Meilenstein soll für die Einheiten der STAGRE die Möglichkeit evaluiert werden, Modellvorhersagen in Analogie zum Umsatz auch für den Bruttoproduktionswert, die Bruttowertschöpfung sowie die Vorleistungen zu berechnen. Da diese Variablen ebenfalls in der Wertschöpfungsstatistik erhoben werden, sind dafür im besten Fall keine tiefgreifenden Anpassungen des Modells erforderlich – vorausgesetzt, die Korrelation dieser Grössen mit den Hilfsvariablen MWST-Umsatz und AHV-Einkommen erweist sich aus ausreichend stark. Diese drei Variablen bilden die Grundlage des Produktionskontos der Volkswirtschaftlichen Gesamtrechnung (VGR) und sind folglich essentiell zur Berechnung des Bruttoinlandsprodukts BIP gemäss Produktionsansatz. Somit könnten die Möglichkeiten für makroökonomische Analysen, etwa zur Bedeutung der multinationalen Unternehmensgruppen für die wirtschaftliche Produktion in der Schweiz, erheblich ausgeweitet werden. Weiterhin wäre der Einbezug weiterer modellgestützter Variablen, namentlich im Bereich der Forschung und Entwicklung (F+E) der Privatwirtschaft, wünschenswert und grundsätzlich denkbar. Dies würde allerdings weitergehende Anpassungen der Modellarchitektur mit sich bringen.

Es wäre zudem erstrebenswert, Varianzschätzer für den hier vorgestellten Modellierungsansatz zu entwickeln, um die statistische Präzision verschiedener Totale für den Umsatz (oder für weitere künftig zu produzierende Variablen) konkret beziffern zu können. Damit stünde ein objektives Kriterium zur Verfügung, auf Grund dessen sich einschätzen liesse, ob in Bezug auf Anforderungen der statistischen Zuverlässigkeit die Publikation von Resultaten nach bestimmten Gliederungsstufen als vertretbar erachtet werden kann oder nicht.

⁵⁰ Die Tabellen aller veröffentlichten Resultate des Umsatzes der STAGRE lassen sich ausgehend von der folgenden Einstiegsseite finden:

<https://www.bfs.admin.ch/bfs/de/home/statistiken/industrie-dienstleistungen/stagre.html>

5 Anhang

5.1 Detailmodelle: Ausführliche Erläuterungen

Der Einbezug von Daten ausserhalb des WS-Rahmens führt zu präziseren Schätzungen der Umsätze. Dies zeigt sich darin, dass für die Modellvariante **dms_onlyinc_s23** (die, in Analogie zu allen Modellen der ersten Etappe, auf den Einbezug von MWST-Umsätzen verzichtet; Resultate siehe Tabelle TA1) ungefähr doppelt so grosse Werte für QM-Global resultieren als für **dms_s23** (der einfachsten Variante, die MWST-Umsätze und Arbeitseinkommen simultan verwendet). Ein Vergleich zwischen **dms_s23** und **dms_onlyvat_s23** offenbart zudem, dass es nicht zielführend wäre, für Beobachtungen mit vorhandenen MWST-Umsätzen auf die in den Arbeitseinkommen vorhandenen Information zu

verzichten, da dies ebenfalls deutlich höhere Werte für QM-Global zur Folge hätte.

In Bezug auf das «Baseline»-Modell **dms_s23** zeigen die folgenden Modellvarianten keine Verbesserung der Prognoseeigenschaften, wenn ausschliesslich QM-Global gewichtet mittels Hochrechnungsgewicht und Branchenfaktoren betrachtet wird: **dms_fte_s23**, **dms_s23a** und **dms_oc_s23**. Die Variante **dms_fte_s23** zeigt zwar Verbesserungen gegenüber **dms_s23**, wenn QM-Global ungewichtet betrachtet wird; jedoch ist das Ausmass dieser Verbesserungen begrenzt, und auch aus theoretischen Erwägungen ist der Baseline-Variante (Verwendung der Arbeitseinkommen statt der Vollzeitäquivalente als Hilfsvariable) den Vorzug zu geben.

Tabelle TA1: Resultate der Evaluation der Detailmodelle

Modellname	QM-Global (RMSE der Abweichungen zwischen modellbasierten Schätzungen und erhobenen Werten, in CHF Mio., N = 99'187)			QM-Aggregat (z-Werte der auf Ebene noga2r aggregierten Schätzungen, M = 419)	
	Ungewichtet	Gew. (HR)	Gew. (HR, Bra.)	Mittelwert	RMS
dms_s23	615.215	49.318	35.497	-0.130	1.829
dms_onlyvat_s23	617.849	52.447	39.931	0.819	3.081
dms_onlyinc_s23	1599.559	112.094	62.238	-0.203	1.151
dms_p_s23	611.829	49.179	35.604	-0.115	1.832
dms_fte_s23	614.139	49.534	36.143	-0.116	1.945
dms_s23a	650.316	51.886	36.201	-0.017	1.955
dms_oc_s23	631.649	52.859	37.107	0.060	2.089
dms_r0_s23	614.297	49.447	35.005	-0.067	1.892
dms_r10_s23	600.129	46.635	31.214	-0.003	1.902
dms_r10_0cs_s23	574.845	46.733	32.800	0.005	1.974
dms_r20_0cs_s23	600.444	50.849	36.411	-0.051	1.992
dms_r30_0cs_s23	576.755	46.885	33.109	-0.031	2.018
dms_r10_0cap_s23	600.601	46.652	31.181	-0.018	1.870
dms_r10_0cs_0cap_s23	575.491	46.750	32.718	-0.010	1.942
dms_r10_0cap_1x_s23	600.765	46.663	31.184	-0.001	1.890
dms_logplus_r10_s23	602.034	46.952	31.447	0.016	1.935
dms_r10_46cs_0cap_s23	569.580	44.900	30.849	-0.017	1.872
dms_r10_46cs_0cap_i_s23	569.569	44.891	30.820	-0.106	1.624

Tabelle TA2: Resultate der Evaluation der Detailmodelle nach Referenzjahren (nur QM-Global gewichtet)

Modellname	QM-Global (RMSE der Abweichungen zwischen modellbasierten Schätzungen und erhobenen Werten, in CHF Mio.), Gewichtet: HR & Branchen						
	2014	2015	2016	2017	2018	2019	2020
dms_s23	37.509	28.688	35.986	32.053	36.432	37.055	38.971
dms_0onlyvat_s23	42.499	33.584	40.043	36.120	41.276	42.659	41.703
dms_onlyinc_s23	66.054	53.950	57.006	59.478	68.457	69.047	57.971
dms_p_s23	37.265	28.674	36.099	32.371	36.864	37.701	38.486
dms_fte_s23	37.799	29.793	36.248	32.814	36.796	38.285	39.633
dms_s23a	37.741	29.585	37.671	32.653	37.364	36.949	39.830
dms_oc_s23	37.151	32.014	36.491	34.382	38.769	39.386	40.215
dms_r0_s23	37.452	28.619	35.887	32.388	36.621	38.422	34.027
dms_r10_s23	28.333	26.145	30.887	28.801	34.742	35.709	32.100
dms_r10_0cs_s23	29.429	28.263	31.361	30.084	37.800	37.599	32.920
dms_r20_0cs_s23	34.022	28.364	34.535	31.383	42.484	41.682	38.836
dms_r30_0cs_s23	29.633	28.827	31.554	30.054	38.061	38.030	33.478
dms_r10_0cap_s23	28.349	26.158	30.905	28.755	34.715	35.535	32.104
dms_r10_0cs_0cap_s23	29.439	28.266	31.373	30.078	37.563	37.193	33.080
dms_r10_0cap_1x_s23	28.349	26.157	30.905	28.751	34.735	35.535	32.105
dms_logplus_r10_s23	28.559	26.288	30.956	29.314	35.016	36.314	31.839
dms_r10_46cs_0cap_s23	28.098	26.119	30.519	28.546	34.005	35.143	31.911
dms_r10_46cs_0cap_i_s23	28.098	26.119	30.519	28.334	34.005	35.143	31.911

Der Einbezug der binären erklärenden Variablen lässt sich nur schwach rechtfertigen, wenn die Präzision der Vorhersage als Mass genommen wird: Wird auf diese Variablen verzichtet (Modellvariante **dms_p_s23**), ist gegenüber **dms_s23** bei QM-Global nur in einer dessen Varianten (mit Hochrechnungsgewichten und Branchenfaktoren) eine geringfügige Verschlechterung zu sehen, während bei den anderen Gewichtungsvarianten eine leichte Verbesserung resultiert. Es wird dennoch vorgeschlagen, am Einbezug dieser binären Variablen festzuhalten: Deren Koeffizienten sind – für jede der drei Variablen und in allen Referenzjahren und beiden Teilmodellen – jeweils statistisch signifikant bei einer Schwelle von 0.1% und haben ein positives Vorzeichen (was plausibel ist). Ein Weglassen dieser binären Variablen würde somit das Risiko mit sich bringen, strukturelle Eigenheiten von Einheiten in multinationalen Gruppen zu unterschlagen, was insbesondere im Kontext der STAGRE problematisch wäre.

Die «Mixed»-Modellvarianten (also jene, die «Random»-Effekte beinhalten) führen zu Verbesserungen der Präzision. Im Falle der einfachsten evaluierten «Mixed»-Variante **dms_r0_s23** fallen

diese bescheiden aus. Die Spezifikation verfeinerter «Random»-Effekte, wie in **dms_r10_s23** führt zu einem deutlicheren und eindeutigen (alle Gewichtungsvarianten von QM-Global umfassenden) Präzisionsgewinn. Wird gegenüber dieser Variante zusätzlich eine Bereinigung unplausibler Koeffizienten-Schätzwerte, wie im Technischen Anhang beschrieben, implementiert (Modellvariante **dms_r10_0cap_s23**), so zeigt sich in QM-Global ungewichtet und gewichtet nur mittels Hochrechnungsgewicht eine geringfügige Verschlechterung der Eigenschaften.⁵¹ Da es für die Bereinigung unplausibler Koeffizienten-Schätzwerte eine gewisse theoretische Rechtfertigung gibt, erscheint eine solche dennoch angezeigt. Schliesslich ergibt sich aus den Werten für QM-Global der verbleibenden «Mixed»-Spezifikationen (**dms_r20_0cs_s23** und **dms_r30_0cs_s23**; diesen Varianten muss für einen kohärenten Vergleich die Variante **dms_r10_0cs_s23** gegenübergestellt werden), dass die «Random»-Effekt-Variante **r10** zu bevorzugen ist.

Die bis hierhin betrachteten Modelle stützen sich alle auf den «hybrid separat/gepoolten» Ansatz. Wird dieser (in der konkreten

⁵¹ Das Modell **dms_r10_s23** ergab, je nach Referenzjahr, für 1500 bis 11 600 Einheiten des Rahmens der WS einen unplausibel tiefen (kleiner als null) geschätzten Wert für γ ; folglich wurden in **dms_r10_0cap_s23** für diese Fälle die Schätzwerte eines reduzierten Alternativmodells mit unterstelltem $\gamma = 0$ verwendet.

Bezogen auf den gesamten WS-Rahmen macht dies zwischen 1% und 7% der Einheiten aus. Fälle mit unplausibel tiefen (kleiner als null) geschätzten Werten für β , für welche spiegelbildlich vorgegangen wurde, waren seltener (zwischen null und 160 Einheiten).

Form des Modells **dms_r10_0cap_s23**) einem Modell mit isolierter Schätzung für die Subpopulation der Unternehmen ohne MWST-Umsätze, aber ansonsten gleicher Spezifikation (**dms_r10_0cap_1x_s23**) gegenübergestellt, sind die Auswirkungen sehr geringfügig: QM-Global verschlechtert sich, unabhängig von der betrachteten Gewichtungsvariante, nur sehr leicht. QM-Global bietet also eine quantitativ nur geringfügige Rechtfertigung für die Verwendung des «hybrid separat/gepoolten» Ansatzes. Aus theoretischer Sicht erscheint es jedoch wünschenswert, die Modellvorhersagen auf eine möglichst umfangreiche Menge an Beobachtungen zu stützen, auf Grund derer das Modell geschätzt werden kann (und der «hybrid separat/gepoolte» Ansatz tut genau dies, indem er die Daten der Einheiten mit MWST-Umsätzen auch in die Schätzung des Modells für die Einheiten ohne MWST-Umsätze einfließen lässt). Somit wird weiterhin **dms_r10_0cap_s23** bevorzugt.

Verworfen wurde, da sich leichte Verschlechterungen in QM-Global in Bezug auf das Vergleichsmodell zeigten, auch die Variante mit Erhöhung der Werte der abhängigen Variablen vor der Logarithmierung (**dms_logplus_r10_s23**).

Eine Betrachtung nach Referenzjahren (Tabelle TA2) wiederum brachte die folgenden Erkenntnisse:

- Mit Ausnahme der Referenzjahre 2017, 2018 und 2019 schneidet **dms_r10_s23** besser oder gleich gut ab wie **dms_r10_0cap_s23**. Da der Vorsprung nie grösser als ein Zehntelprozent ist, und da sich die Bereinigung unplausibler Koeffizienten-Schätzwerte durch theoretische Überlegungen rechtfertigen lässt, scheint die Bevorzugung von **dms_r10_0cap_s23** weiterhin angezeigt.
- Das Modell **dms_r10_0cap_1x_s23** (Verzicht auf den «hybrid separat/gepoolten» Ansatz) führt 2015 und 2017 zu einer leichten Verringerung von QM-Global gegenüber dem Vergleichsmodell (**dms_r10_0cap_s23**). Da in den anderen Referenzjahren keine Verbesserung feststellbar ist, und insbesondere gestützt auf die theoretischen Erwägungen weiter oben, erscheint ein Verzicht auf den «hybrid separat/gepoolten» Ansatz dennoch nicht angezeigt.

Tabelle TA3: Resultate der Evaluation der Detailmodelle nach NOGA-Abschnitten (nur QM-Global gewichtet)

Modellname	QM-Global (RMSE der Abweichungen zwischen modellbasierten Schätzungen und erhobenen Werten, in CHF Mio.), Gewichtet: HR & Branchen															
	B	C	D	E	F	G	H	I	J	L	M	N	P	Q	R	S
dms_s23	7.033	47.949	551.261	4.421	6.347	120.605	108.560	2.320	32.596	6.036	39.500	21.136	5.471	1.445	19.196	3.127
dms_onlyvat_s23	17.115	61.176	706.461	3.831	4.960	126.276	116.763	3.266	30.405	5.938	38.368	25.718	5.291	3.078	16.596	3.406
dms_onlyinc_s23	11.804	62.002	462.621	13.336	15.309	331.695	256.742	2.909	62.463	6.030	43.779	28.157	6.743	1.280	25.736	3.516
dms_p_s23	7.119	49.775	544.461	4.464	6.166	120.020	105.125	2.298	32.697	6.026	39.446	21.321	5.484	1.464	19.100	3.212
dms_fte_s23	6.502	50.298	548.833	4.278	6.487	120.773	109.640	2.484	31.379	6.245	39.790	25.157	5.776	1.427	18.758	3.250
dms_s23a	8.587	48.759	561.503	4.263	8.378	131.750	107.950	2.575	33.737	6.484	38.983	21.236	5.342	1.336	18.483	3.158
dms_oc_s23	6.542	50.373	759.379	4.376	4.648	127.899	86.028	2.427	31.034	6.508	38.614	20.003	5.620	1.163	19.027	2.777
dms_r0_s23	7.129	47.486	584.255	4.334	5.806	120.314	94.431	2.321	30.845	6.150	39.616	20.992	5.439	1.305	19.226	2.944
dms_r10_s23	2.540	45.171	496.103	4.103	5.472	113.088	39.946	1.962	27.904	5.890	38.190	18.595	4.387	1.297	18.676	2.492
dms_r10_0cs_s23	2.681	46.275	633.894	3.982	5.933	107.733	50.519	2.214	28.824	7.977	38.463	18.940	4.554	1.255	18.640	2.225
dms_r20_0cs_s23	2.831	48.294	846.256	3.898	5.246	112.379	65.141	2.281	32.068	7.934	38.842	18.749	5.098	1.178	19.129	2.619
dms_r30_0cs_s23	2.870	46.499	629.638	3.895	5.974	108.386	60.279	2.218	29.716	8.044	38.461	18.853	4.664	1.177	18.589	2.285
dms_r10_0cap_s23	2.537	45.154	496.100	4.135	5.476	113.202	38.467	1.914	27.740	5.890	38.191	18.586	4.388	1.297	18.676	2.492
dms_r10_0cs_0cap_s23	2.668	46.256	633.889	4.004	5.955	107.883	47.848	2.186	28.239	7.977	38.464	18.943	4.556	1.255	18.640	2.225
dms_r10_0cap_1x_s23	2.509	45.150	496.087	4.136	5.496	113.231	38.442	1.915	27.560	5.961	38.244	18.584	4.385	1.321	18.676	2.489
dms_logplus_r10_s23	2.616	45.639	509.991	4.084	5.524	113.388	35.716	2.088	28.317	6.049	38.428	18.882	4.608	1.287	18.661	2.528
dms_r10_46cs_0cap_s23	2.537	45.154	496.100	4.135	5.476	106.646	38.467	1.914	27.740	5.890	38.191	18.586	4.388	1.297	18.676	2.492
dms_r10_46cs_0cap_i_s23	2.537	45.154	496.100	4.135	5.476	106.646	38.467	1.914	27.740	5.890	38.191	17.921	4.388	1.297	18.676	2.492

Tabelle TA4: Abschnitte und Abteilungen der NOGA

Ab-schn.	Abt.	Bezeichnung
A	01	Landwirtschaft, Jagd und damit verbundene Tätigkeiten
	02	Forstwirtschaft und Holzeinschlag
	03	Fischerei und Aquakultur
B	05	Kohlenbergbau
	06	Gewinnung von Erdöl und Erdgas
	07	Erzbergbau
	08	Gewinnung von Steinen und Erden, sonstiger Bergbau
	09	Erbringung von Dienstleistungen für den Bergbau und für die Gewinnung von Steinen und Erden
C	10	Herstellung von Nahrungs- und Futtermitteln
	11	Getränkeherstellung
	12	Tabakverarbeitung
	13	Herstellung von Textilien
	14	Herstellung von Bekleidung
	15	Herstellung von Leder, Lederwaren und Schuhen
	16	Herstellung von Holz-, Flecht-, Korb- und Korkwaren (ohne Möbel)
	17	Herstellung von Papier, Pappe und Waren daraus
	18	Herstellung von Druckerzeugnissen; Vervielfältigung von bespielten Ton-, Bild- und Datenträgern
	19	Kokerei und Mineralölverarbeitung
	20	Herstellung von chemischen Erzeugnissen
	21	Herstellung von pharmazeutischen Erzeugnissen
	22	Herstellung von Gummi- und Kunststoffwaren
	23	Herstellung von Glas und Glaswaren, Keramik, Verarbeitung von Steinen und Erden
	24	Metallerzeugung und -bearbeitung
	25	Herstellung von Metallerzeugnissen
	26	Herstellung von Datenverarbeitungsgeräten, elektronischen und optischen Erzeugnissen
	27	Herstellung von elektrischen Ausrüstungen
	28	Maschinenbau
	29	Herstellung von Automobilen und Automobilteilen
	30	Sonstiger Fahrzeugbau
	31	Herstellung von Möbeln
	32	Herstellung von sonstigen Waren
	33	Reparatur und Installation von Maschinen und Ausrüstungen
D	35	Energieversorgung
E	36	Wasserversorgung
	37	Abwasserentsorgung
	38	Sammlung, Behandlung und Beseitigung von Abfällen; Rückgewinnung
	39	Beseitigung von Umweltverschmutzungen und sonstige Entsorgung
F	41	Hochbau
	42	Tiefbau
	43	Vorbereitende Baustellenarbeiten, Bauinstallation und sonstiges Ausbaugewerbe
G	45	Handel mit Motorfahrzeugen; Instandhaltung und Reparatur von Motorfahrzeugen
	46	Grosshandel (ohne Handel mit Motorfahrzeugen)
	47	Detailhandel (ohne Handel mit Motorfahrzeugen)

Die durch gestrichelte Linien getrennten Abteilungen bezeichnen jene Abteilungsgruppen, die für die Konstruktion der für die Modellierung verwendeten Gliederungsstufe **noga2r** zusammengelegt wurden (05–09, 11–12, 19–20, 38–39).

Ab-schn.	Abt.	Bezeichnung
H	49	Landverkehr und Transport in Rohrfernleitungen
	50	Schifffahrt
	51	Luftfahrt
	52	Lagerei sowie Erbringung von sonstigen Dienstleistungen für den Verkehr
	53	Post-, Kurier- und Expressdienste
I	55	Beherbergung
	56	Gastronomie
J	58	Verlagswesen
	59	Herstellung, Verleih und Vertrieb von Filmen und Fernsehprogrammen; Kinos; Tonstudios und Verlegen von Musik
	60	Rundfunkveranstalter
	61	Telekommunikation
	62	Erbringung von Dienstleistungen der Informationstechnologie
	63	Informationsdienstleistungen
K	64	Erbringung von Finanzdienstleistungen
	65	Versicherungen, Rückversicherungen und Pensionskassen (ohne Sozialversicherung)
	66	Mit Finanz- und Versicherungsdienstleistungen verbundene Tätigkeiten
L	68	Grundstücks- und Wohnungswesen
M	69	Rechts- und Steuerberatung, Wirtschaftsprüfung
	70	Verwaltung und Führung von Unternehmen und Betrieben; Unternehmensberatung
	71	Architektur- und Ingenieurbüros; technische, physikalische und chemische Untersuchung
	72	Forschung und Entwicklung
	73	Werbung und Marktforschung
	74	Sonstige freiberufliche, wissenschaftliche und technische Tätigkeiten
	75	Veterinärwesen
N	77	Vermietung von beweglichen Sachen
	78	Vermittlung und Überlassung von Arbeitskräften
	79	Reisebüros, Reiseveranstalter und Erbringung sonstiger Reservierungsdienstleistungen
	80	Wach- und Sicherheitsdienste sowie Detekteien
	81	Gebäudebetreuung; Garten- und Landschaftsbau
	82	Erbringung von wirtschaftlichen Dienstleistungen für Unternehmen und Privatpersonen a. n. g.
O	84	Öffentliche Verwaltung, Verteidigung; Sozialversicherung
P	85	Erziehung und Unterricht
Q	86	Gesundheitswesen
	87	Heime (ohne Erholungs- und Ferienheime)
	88	Sozialwesen (ohne Heime)
R	90	Kreative, künstlerische und unterhaltende Tätigkeiten
	91	Bibliotheken, Archive, Museen, botanische und zoologische Gärten
	92	Spiel-, Wett- und Lotteriewesen
	93	Erbringung von Dienstleistungen des Sports, der Unterhaltung und der Erholung
S	94	Interessenvertretungen sowie kirchliche und sonstige religiöse Vereinigungen (ohne Sozialwesen und Sport)
	95	Reparatur von Datenverarbeitungsgeräten und Gebrauchsgütern
	96	Erbringung von sonstigen überwiegend persönlichen Dienstleistungen
T	97	Private Haushalte mit Hauspersonal
	98	Herstellung von Waren und Erbringung von Dienstleistungen durch private Haushalte für den Eigenbedarf ohne ausgeprägten Schwerpunkt
U	99	Exterritoriale Organisationen und Körperschaften

6 Technischer Anhang

Grundlagen der Spezifikation für die Modellierung des Umsatzes

Der Zweck des Modells ist, für die Einheiten i der Grundgesamtheit aller Unternehmen (oder einer Untermenge davon, im konkreten Fall aller Unternehmen des STAGRE-Universums) eine möglichst präzise Vorhersage zu deren Umsatz U_i (also dem von der Wertschöpfungsstatistik WS gemessen “Umsatz betriebswirtschaftlich”) zu liefern, basierend auf den Hilfsvariablen “Umsatz gemäss MWST-Statistik” T_i und “Arbeitseinkommen gemäss AHV” L_i .¹

Unterstellt wird der folgende multiplikative Zusammenhang:

$$U_i = A_{g(i)} T_i^{\beta_{g(i)}} L_i^{\gamma_{g(i)}} e^{\varepsilon_i} \quad (1)$$

Wobei $g(i)$ die Zugehörigkeit von i zu bestimmten Subgruppen (etwa NOGA-Abteilungen oder -Arten) des Universums ausdrückt, zwischen denen eine gewisse Heterogenität in Bezug auf die Parameter A , β und γ vermutet wird, welcher durch die Modellierung Rechnung zu tragen ist. Für die nun folgenden Erläuterungen ist vorerst nicht relevant, welche Gliederungsstufe der NOGA zur Bildung dieser Subgruppen am sinnvollsten ist, und ob allenfalls noch weitere Kriterien zur Bildung von Subgruppen einzubeziehen sind.

Es wird also ein exponentieller oder linear homogener (letzteres als Spezialfall, wenn die Gesamtelastizität $\beta_{g(i)} + \gamma_{g(i)} = 1$ beträgt) Zusammenhang unterstellt zwischen dem zu schätzenden WS-Umsatz und dem Produkt von MWST-Umsatz und AHV-Arbeitseinkommen (beide jeweils exponiert mit den empirisch zu bestimmenden Exponenten $\beta_{g(i)}$ und $\gamma_{g(i)}$).² Die Variabilität des Anpassungsfaktors A in Bezug auf $g(i)$ trägt der Tatsache Rechnung, dass $E(U_i|T_i, L_i)$ in Abhängigkeit von $g(i)$ variieren kann; das heisst, für gegebene Werte von MWST-Umsatz und AHV-Arbeitseinkommen kann der erwartete WS-Umsatz je nach NOGA-Branchen

etc. unterschiedlich hoch ausfallen. Der Zufallsterm e^{ε_i} absorbiert sämtliche in der Modellierung nicht berücksichtigten Faktoren.

Die gewählte Spezifikation hat die Eigenschaft, dass sich zwischen den Logarithmen der relevanten Variablen ein linearer Zusammenhang ergibt:

$$u_i = \alpha_{g(i)} + \beta_{g(i)} t_i + \gamma_{g(i)} l_i + \varepsilon_i \quad (2)$$

Wobei hier und im Folgenden, zwecks einfacherer Darstellung, die Logarithmen der Variablen mittels derer Kleinbuchstaben ausgedrückt werden, also beispielsweise $\log U_i = u_i$, und die Logarithmen der gruppenspezifischen Anpassungsfaktoren A mittels α .³

Unter bestimmten Nichtkorrelationsannahmen bezüglich ε_i kann (2) für die statistische Modellierung des Erwartungswerts des Logarithmus von U_i verwendet werden. Insbesondere darf ε_i einerseits weder mit t_i noch mit l_i und andererseits nicht mit der Gegebenheit, ob i Teil der WS-Nettostichprobe ist oder nicht, korrelieren. Zur Operationalisierung dieses Modells sind die folgenden Präzisierungen angebracht:

- Wie schon erläutert können die Parameter α , β und γ zwischen den Subgruppen $g(i)$ variieren. Damit wird insbesondere auch der Tatsache Rechnung getragen, dass der MWST-Umsatz in gewissen Branchen ein relativ zuverlässiger Proxy für den WS-Umsatz darstellt. In anderen Branchen – insbesondere solche, deren Produktion mehrheitlich von der MWST ausgenommen ist – dürfte dies weniger der Fall sein, hingegen könnten hier die Arbeitseinkommen gemäss AHV aussagekräftiger sein, was sich in tieferen $\beta_{g(i)}$ und analog dazu höheren $\gamma_{g(i)}$ niederschlagen würde.
- Die soeben erwähnte Variabilität lässt sich – entsprechend der klassischen OLS-Annahmen – in der Form von Fix-Effekten bzw. fixen Interaktionen spezifizieren, oder aber im Rahmen von *Linear Mixed Models* (LMM) als Random-Effekte (RE).⁴ Für die nun folgenden Erläuterungen steht

¹Als Alternative zu den Arbeitseinkommen wurde auch die Hilfsvariable “Beschäftigung gemäss STATENT” evaluiert. Im Kontext der vorliegenden Überlegungen wird der Einfachheit halber nur auf die Arbeitseinkommen Bezug genommen; diese Überlegungen sind jedoch ohne weiteres auch auf die alternative Hilfsvariable anwendbar.

²In der Modellspezifikation kann, falls dies als sinnvoll erachtet wird, ohne grösseren Aufwand auch ein linear homogener Zusammenhang erzwungen werden, indem $\gamma_{g(i)} = 1 - \beta_{g(i)}$ unterstellt und die Schätzgleichung wie folgt angepasst wird: $U_i/L_i = A_{g(i)} (T_i/L_i)^{\beta_{g(i)}} e^{\varepsilon_i}$

³Im Falle eines erzwungenen linear homogenen Zusammenhangs:

$$u_i - l_i = \alpha_{g(i)} + \beta_{g(i)} (t_i - l_i) + \varepsilon_i$$

⁴Formell lässt sich in einem LMM beispielsweise der Koeffizient $\alpha_{g(i)}$ als Summe $\alpha_{G[g(i)]}^F + \alpha_{g(i)}^R$ definieren, wobei $\alpha_{G[g(i)]}^F$

die LMM-Variante im Vordergrund, da sie (a) unter bestimmten Voraussetzungen eine detailliertere Gliederung der Subgruppen $g(i)$ und (b) potentiell Effizienzsteigerungen bei der Erweiterung des Modellierungsansatzes auf Einheiten mit fehlendem MWST-Umsatz (siehe nächster Abschnitt) zulässt. Eine Umsetzung mittels OLS/FE ist mit gewissen Einschränkungen jedoch auch denkbar (und zumindest zur Robustheitsprüfung der LMM/RE-Varianten in Betracht zu ziehen).

- Das Modell lässt sich mit relativ geringem Aufwand erweitern, um der Tatsache Rechnung zu tragen, dass für manche Einheiten keine original erhobenen MWST-Umsätze verfügbar sind, sondern imputierte, partiell imputierte oder verteilte MWST-Umsätze; und dass bisherige Analysen gezeigt haben, dass diese nicht originalen MWST-Umsätze statistisch weniger zuverlässig sind als originale. Folglich ist zu erwägen, die Subgruppen $g(i)$ nicht nur auf der Basis der Wirtschaftstätigkeit (NOGA), sondern zusätzlich auch nach der Methode der MWST-Umsatzerhebung zu spezifizieren.
- Um eine Einschätzung zu ermöglichen, ob diese flexible, eher aufwändige Modellanlage gerechtfertigt ist, wurden auch vereinfachte Varianten des Modells getestet, in welchen für eine gegebene Einheit i entweder $\beta_{g(i)} = 0$ oder $\gamma_{g(i)} = 0$ unterstellt wird; das heisst, es wird ausschliesslich jeweils auf den MWST-Umsatz oder auf die AHV-Arbeitseinkommen zurückgegriffen. Die Vorhersageeigenschaften dieser vereinfachten Modelle lassen sich dann mit jenen des flexiblen Modells vergleichen. Eine Implementierung des flexiblen Modells ist dann gerechtfertigt, wenn damit wesentlich bessere Vorhersageeigenschaften resultieren.

Behandlung von Einheiten mit fehlendem MWST-Umsatz

Es gibt Fälle von Einheiten, die zwar einen Umsatz realisieren, jedoch keine MWST deklarieren.⁵ Im Referenzjahr 2016 waren von den 14'000 Einheiten mit Umsatz gemäss WS mehr als 1000 von diesem Phänomen

die als *fixed* und $\alpha_{g(i)}^R$ die als *random* zu modellierende Komponente ist. $G[g(i)]$ beschreibt hier eine Zusammenfassung (also eine Gruppierung mit reduziertem Detaillierungsgrad) von $g(i)$, die aus theoretischen oder empirischen Erwägungen Sinn ergibt, für welche die Effekte als *fixed* zu modellieren sind. Denkbar ist auch $\alpha_{G[g(i)]}^F = \alpha^F$, also dass ein einziger, für die gesamte Population homogener Fixeffekt unterstellt wird. In Analogie zu $\alpha_{g(i)}$ ist es angezeigt, auch $\beta_{g(i)}$ und $\gamma_{g(i)}$ als *mixed effects* zu modellieren.

⁵Es handelt sich um Einheiten, deren Umsatz entweder geringfügig ist (unterhalb der Schwelle von CHF 100'000, ab welcher die Deklaration der MWST obligatorisch ist) oder deren Leistungen von der Steuer ausgenommen sind. Zu den Einheiten ohne MWST-Umsätze werden im vorliegenden Kontext auch jene gezählt, für die die Sektion Unternehmensregisterdaten (URD) des BFS in den MWST-Daten einen Umsatz imputiert hat.

betroffen. Die vorgeschlagene Schätzgleichung erlaubt die Berücksichtigung solcher Einheiten a priori nicht. Für dieses Problem sind verschiedene Lösungsansätze denkbar. Wir werden im folgenden ausschliesslich Lösungsansätze betrachten, in denen für solche Einheiten die Existenz eines latenten positiven MWST-Umsatzes T_i^* postuliert wird, der – eingesetzt als T_i – den in (1) unterstellten Zusammenhang erfüllt, und der in einem log-linearen Zusammenhang mit den (beobachteten) Arbeitseinkommen steht:

$$T_i^* = A_{g(i)}^{tl} L_i^{\delta_{g(i)}^{tl}} e^{\eta_i^{tl}} \quad (3)$$

Dies entspricht, in logarithmischer Formulierung:

$$t_i^* = \alpha_{g(i)}^{tl} + \delta_{g(i)}^{tl} l_i + \eta_i^{tl} \quad (4)$$

Nach Einsetzung in die Basisgleichung (2) resultiert:

$$u_i = \left(\alpha_{g(i)} + \beta_{g(i)} \alpha_{g(i)}^{tl} \right) + \left(\gamma_{g(i)} + \beta_{g(i)} \delta_{g(i)}^{tl} \right) l_i + (\varepsilon_i + \beta_{g(i)} \eta_i^{tl}) \quad (5)$$

Das heisst, der Erwartungswert des logarithmierten WS-Umsatzes für Einheiten mit nicht vorhandenem MWST-Umsatz ist eine lineare Funktion der logarithmierten AHV-Arbeitseinkommen. Er lässt sich aus dieser berechnen, wenn nebst den bereits in (2) postulierten Parametern zusätzlich $\alpha_{g(i)}^{tl}$ und $\delta_{g(i)}^{tl}$ – oder, alternativ dazu, die Terme $\beta_{g(i)} \alpha_{g(i)}^{tl}$ und $\beta_{g(i)} \delta_{g(i)}^{tl}$ – identifiziert werden können.

Der konzeptionell einfachste Lösungsansatz besteht darin, (2) und (5) **separat** für die jeweils betroffenen Subpopulationen (Einheiten mit vorhandenem bzw. fehlendem MWST-Umsatz) zu schätzen, wobei in (5) nicht die einzelnen Komponenten, sondern die Terme $\alpha_{g(i)} + \beta_{g(i)} \alpha_{g(i)}^{tl}$ sowie $\gamma_{g(i)} + \beta_{g(i)} \delta_{g(i)}^{tl}$ identifiziert werden. Für die letztere der beiden Subpopulationen dürfte dies, da die Anzahl der Einheiten bedeutend geringer ist als in ersterer, eher ineffizient sein; insbesondere wenn eine Schätzung der Terme gemäss der selben detaillierten Untergliederung in Subgruppen $g(i)$ wie für die Subpopulation mit MWST-Umsätzen angestrebt wird.

Wird unterstellt, dass der in (3) postulierte Zusammenhang zwischen L_i und T_i^* für Einheiten mit fehlendem MWST-Umsatz derselbe ist wie zwischen L_i und T_i für Einheiten mit vorhandenem MWST-Umsatz, ist eine **sequentielle** Schätzung von (2) und (4) denkbar. Das heisst, beide Gleichungen werden innerhalb der Subpopulation mit vorhandenen MWST-Umsätzen geschätzt, um daraus die Parameter in (5) zu berechnen, einzusetzen und damit u_i für Einheiten mit fehlenden MWST-Umsätzen zu ermitteln. Dieses Vorgehen ist potentiell effizienter (da auf eine bedeutend grössere Anzahl

von Beobachtungen aufbauend), jedoch ist die eingangs erwähnte Unterstellung eines identischen Zusammenhangs für beide Subpopulationen sehr unrealistisch, da fehlende MWST-Umsätze tendenziell eher in Einheiten mit geringfügigem Umsatz auftreten dürften. Formell: $E(u_i | l_i, T_i \text{ fehlt}) < E(u_i | l_i, T_i > 0)$; das heisst, für $(\alpha_{g(i)}^{tl}, \delta_{g(i)}^{tl})$ auf Grund der Subpopulation mit vorhandenen MWST-Umsätzen mittels (4) ermittelte Schätzwerte sind, bezogen auf die Subpopulation mit fehlenden MWST-Umsätzen, mit hoher Wahrscheinlichkeit verzerrt und hätten zu hohe Vorhersagen für u_i zur Folge.⁶ Dieser Lösungsansatz ist somit ebenfalls nicht zielführend.

Die beiden Gleichungen (2) und (5) lassen sich auch **gepoolt** schätzen, also durch Verwendung einer auf geeignete Art konstruierten Datenmatrix, welche die Vereinigungsmenge der beiden Subpopulationen abdeckt. Falls keine Restriktionen oder Korrelationen zwischen den Parametern $\alpha_{g(i)}$, $\beta_{g(i)}$ und $\gamma_{g(i)}$ in (2) einerseits und den Termen in (5) andererseits unterstellt werden, ist dieser Ansatz (zumindest bezüglich der Punktschätzungen) identisch zu separaten Schätzungen der beiden Gleichungen. Werden solche Restriktionen und/oder (im Rahmen eines LMM-Ansatzes) Korrelation unterstellt, ist jedoch denkbar, dass sich die Terme in (5) effizienter schätzen lassen, und dass dadurch präzisere Vorhersagen des WS-Umsatzes für Einheiten ohne MWST-Umsätze resultieren.

Ein verwandter Lösungsansatz besteht in einer **hybrid separat/gepoolten** Schätzung von (2) und (5); konkret einer separaten Schätzung von (2) zwecks Vorhersage des Umsatzes für Einheiten mit vorhandenem MWST-Umsatz, und einer gepoolten Schätzung von (2) und (5) unter Weglassung von t_i zwecks Vorhersage des Umsatzes für Einheiten mit fehlendem Umsatz. Ähnlich der oben erläuterten “rein” gepoolten Schätzung sind hier gegenüber separaten Schätzungen Effizienzgewinne zu erwarten, wobei der wesentliche Unterschied zum rein gepoolten Ansatz darin besteht, dass keine Annahmen zu Restriktionen oder Korrelationen zwischen $\beta_{g(i)}$ einerseits und den Termen $\beta_{g(i)}\alpha_{g(i)}^{tl}$ und $\beta_{g(i)}\delta_{g(i)}^{tl}$ andererseits unterstellt werden können. Es werden also konkret geschätzt:

$$u_i = \alpha_{g(i)} + \beta_{g(i)}t_i + \gamma_{g(i)}l_i + \varepsilon_i \quad \forall i \mid T_i > 0$$

⁶Diese Vermutung lässt sich durch die Resultate der Schätzung von Gleichung (6) weiter unten erhärten, wo die Schätzwerte für die *fixed* Komponente von $\gamma_{g(i)}^{L0}$ signifikant tiefer liegen als jene von $\gamma_{g(i)}^{L1}$, und zwar bei einer Signifikanzschwelle von 0.01 für zwei von fünf Referenzjahren (und von 0.1 für alle Referenzjahre).

$$u_i = \alpha_{g(i)}^{L0}[1 - I(T_i > 0)] + \alpha_{g(i)}^{L1}I(T_i > 0) + \gamma_{g(i)}^{L0}[1 - I(T_i > 0)]l_i + \gamma_{g(i)}^{L1}I(T_i > 0)l_i + (\varepsilon_i + \beta_{g(i)}\eta_i^{tl}) \quad \forall i \quad (6)$$

Wobei $I(T_i > 0)$ eine binäre Funktion ist, die den Wert von 1 liefert falls T_i vorhanden und positiv ist, und 0 andernfalls. Von den Schätzparametern der zweiten Gleichung werden – zur Vorhersage von u_i für Einheiten ohne MWST-Umsätze – lediglich $\alpha_{g(i)}^{L0}$ und $\gamma_{g(i)}^{L0}$ weiterverwendet. Diese beiden Parameter erfassen die Terme $\alpha_{g(i)} + \beta_{g(i)}\alpha_{g(i)}^{tl}$ bzw. $\gamma_{g(i)} + \beta_{g(i)}\delta_{g(i)}^{tl}$ aus (5). Die differenzierte Spezifikation (Unterscheidung zwischen $\alpha_{g(i)}^{L0}$ und $\alpha_{g(i)}^{L1}$ einerseits, und $\gamma_{g(i)}^{L0}$ und $\gamma_{g(i)}^{L1}$ andererseits) trägt der weiter oben erwähnten Vermutung Rechnung, dass es unzulässig wäre, die Parameter des in (3) unterstellten Zusammenhangs für die Subpopulation der Einheiten ohne MWST-Umsatz einfach mittels der Subpopulation der Einheiten mit MWST-Umsatz zu schätzen. Gleichzeitig erlaubt sie, indem sie Korrelation zwischen den als Random-Effekten modellierten Parametern $\alpha_{g(i)}^{L0}$, $\alpha_{g(i)}^{L1}$, $\gamma_{g(i)}^{L0}$ und $\gamma_{g(i)}^{L1}$ zulässt,⁷ die zahlreicher vorhandenen Daten der Einheiten mit vorhandenen MWST-Umsätzen in die Schätzung der Parameter für Einheiten mit fehlenden MWST-Umsätzen einfließen zu lassen, wodurch Effizienzgewinne möglich sind.

Für die Operationalisierung des Modells wurde schliesslich dem hybrid separat/gepoolten Ansatz den Vorzug zu geben. Dies, da bei den getesteten LMM-Schätzungen sehr schnell Konvergenzprobleme auftraten, wenn möglichst flexible (bezüglich der Definition der Subgruppen $g(i)$ sowie der Korrelationsstruktur zwischen den Random-Effekten) Formulierungen für gepoolte Schätzungen gewählt wurden. Es drängen sich also Vereinfachungen bezüglich der Korrelationsannahmen der Parameter auf. Die Nichtberücksichtigung allfälliger Korrelationen zwischen $\beta_{g(i)}\alpha_{g(i)}^{tl}$ und $\beta_{g(i)}\delta_{g(i)}^{tl}$ einerseits und $\beta_{g(i)}$ (im “rein” gepoolten Ansatz könnten solche theoretisch berücksichtigt werden) scheint vertretbar; insbesondere auch weil sich die quantitativen Unterschiede in den Resultaten gegenüber den punktuell getesteten Varianten rein gepoolter Ansätze in Grenzen halten.

Behandlung von Subgruppen mit unplausiblen Schätzwerten der Koeffizienten

Der LMM-Ansatz bringt ein grosses Potential bezüglich

⁷Im Rahmen des Schätzverfahrens wird hierzu eine Kovarianzmatrix Σ mit Dimension $k \times k$ berechnet, wobei $k = 4$ (die Anzahl der Parameter). Dies impliziert (da Σ symmetrisch ist) die Schätzung von $(k - 1)k/2 = 6$ zusätzlichen Korrelationsparametern gegenüber einer Variante ohne Berücksichtigung von Korrelationen zwischen den Random-Effekten. Siehe Bates, D., Mächler, M., Bolker, B., & Walker, S. (2015). Fitting Linear Mixed-Effects Models Using lme4. Journal of Statistical Software, 67(1), 1–48. <https://doi.org/10.18637/jss.v067.i01>

Modellpräzision mit sich. Es hat sich jedoch herausgestellt, dass in einigen Subgruppen $g(i)$ unplausible Schätzwerte der Koeffizienten $\beta_{g(i)}$ oder $\gamma_{g(i)}$ resultieren können. Als unplausibel zu betrachten sind in erster Linie Werte kleiner als null, da diese zur Folge haben, dass der vorhergesagte Wert für den WS-Umsatz bei einer Erhöhung des Werts der betroffenen Hilfsvariablen (unter Beibehaltung der Werte für die anderen Parameter) sinkt.⁸ In einem Modell mit erzwungener linearer Homogenität (also mit $\beta_{g(i)} + \gamma_{g(i)} = 1$) gilt zudem $\hat{\beta}_{g(i)} < 0 \Leftrightarrow \hat{\gamma}_{g(i)} > 1$ bzw. $\hat{\gamma}_{g(i)} < 0 \Leftrightarrow \hat{\beta}_{g(i)} > 1$; das heisst, Schätzwerte kleiner null für den einen Koeffizienten gehen einher mit Schätzwerten grösser eins für den anderen Koeffizienten. Während Koeffizientenwerte grösser eins nicht per se ökonomisch unplausibel sind (diese können beispielsweise Ausdruck von Skaleneffekten sein, wie sie in der Realität unter bestimmten Umständen auftreten), bergen sie ein gewisses Risiko einer Überschätzung einzelner Werte der zu modellierenden Variablen, die insbesondere in Gegenwart von ausgeprägt rechtsschief verteilten Variablen (wie sie der WS-Umsatz und dessen Hilfsvariablen sind) unverhältnismässig stark auf die aggregierten Werte durchschlagen können.

Um dieses Problem abzumildern, werden deshalb Modelle vorgeschlagen, die nach dem folgenden Vorgehen aufgebaut sind:

1. Schätzen des Ausgangsmodells, also der Gleichung (2) und Berechnen sowohl der vorhergesagten Werte $\hat{u}_i$ als auch der Schätzwerte $\hat{\beta}_{g(i)}$ und $\hat{\gamma}_{g(i)}$ für alle Einheiten i der Grundgesamtheit;
2. Bestimmung der Menge $G = \{i | \hat{\beta}_{g(i)} < 0\}$, also aller Einheiten der Grundgesamtheit, für welche der für die Vorhersage von u_i verwendete Schätzkoeffizient für $\beta_{g(i)}$ im Ausgangsmodell kleiner als null ist;
3. Ersatz, für alle $i \in G$, von $\hat{u}_i$ durch die Vorhersage $\hat{u}_i^{L1}$, definiert als $\hat{u}_i^{L1} = \hat{\alpha}_{g(i)}^{L1} + \hat{\gamma}_{g(i)}^{L1} l_i$; also der Vorhersage aus dem gemäss (6) geschätzten Modell;⁹
4. Anwenden der beiden vorherigen Schritte, in spiegelbildlicher Art und Weise, für die Menge $B = \{i | \hat{\gamma}_{g(i)} < 0\}$, wobei hier für den Ersatz von $\hat{u}_i$ auf die Vorhersage gemäss einem vereinfachten Modell der Spezifikation $u_i = \alpha_{g(i)}^T + \beta_{g(i)}^T t_i + \varepsilon_i^T$ zurückgegriffen wird.

⁸Der LMM-Ansatz ist gegenüber einem FE-Ansatz deshalb anfälliger auf dieses Phänomen, weil er einerseits eine feinere Definition der Subgruppen $g(i)$ und andererseits die Spezifikation von verschachtelten ("nested") Random-Effekten ermöglicht. Im bevorzugten Modell `dms_r10` resultieren, je nach Referenzjahr, für zwischen 1% und 7% der Beobachtungen ein $\hat{\gamma}_{g(i)} < 0$, und für zwischen 0 und 0.1% der Beobachtungen ein $\hat{\beta}_{g(i)} < 0$.

⁹Somit wird der Tatsache Rechnung getragen, dass ein MWST-Umsatz zwar vorhanden ist, dessen Einbezug in die Vorhersage aber nicht als sinnvoll erachtet wird.

Falls für das Ausgangsmodell die Variante mit erzwungener linearer Homogenität gewählt wurde, ist es aus Konsistenzgründen sinnvoll, auch in den in Punkt (3) herangezogenen "Ausweichmodellen" lineare Homogenität zu unterstellen (das heisst, gemäss der oben verwendeten Darstellungsweise, auf die Schätzgleichung $u_i^{L1} - l_i = \alpha_{g(i)}^{L1} + \varepsilon_i$ zurückzugreifen).

Publikationsprogramm BFS

Das Bundesamt für Statistik (BFS) hat als zentrale Statistikstelle des Bundes die Aufgabe, statistische Informationen zur Schweiz breiten Benutzerkreisen zur Verfügung zu stellen. Die Verbreitung geschieht gegliedert nach Themenbereichen und mit verschiedenen Informationsmitteln über mehrere Kanäle.

Die statistischen Themenbereiche

- 00 Statistische Grundlagen und Übersichten
- 01 Bevölkerung
- 02 Raum und Umwelt
- 03 Arbeit und Erwerb
- 04 Volkswirtschaft
- 05 Preise
- 06 Industrie und Dienstleistungen
- 07 Land- und Forstwirtschaft
- 08 Energie
- 09 Bau- und Wohnungswesen
- 10 Tourismus
- 11 Mobilität und Verkehr
- 12 Geld, Banken, Versicherungen
- 13 Soziale Sicherheit
- 14 Gesundheit
- 15 Bildung und Wissenschaft
- 16 Kultur, Medien, Informationsgesellschaft, Sport
- 17 Politik
- 18 Öffentliche Verwaltung und Finanzen
- 19 Kriminalität und Strafrecht
- 20 Wirtschaftliche und soziale Situation der Bevölkerung
- 21 Nachhaltige Entwicklung, regionale und internationale Disparitäten

Die zentralen Übersichtspublikationen

Statistisches Jahrbuch der Schweiz



Das vom Bundesamt für Statistik (BFS) herausgegebene Statistische Jahrbuch ist seit 1891 das Standardwerk der Schweizer Statistik. Es fasst die wichtigsten statistischen Ergebnisse zu Bevölkerung, Gesellschaft, Staat, Wirtschaft und Umwelt des Landes zusammen.

Taschenstatistik der Schweiz



Die Taschenstatistik ist eine attraktive, kurzweilige Zusammenfassung der wichtigsten Zahlen eines Jahres. Die Publikation mit 52 Seiten im praktischen A6/5-Format ist gratis und in fünf Sprachen (Deutsch, Französisch, Italienisch, Rätoromanisch und Englisch) erhältlich.

Das BFS im Internet – www.statistik.ch

Das Portal «Statistik Schweiz» bietet Ihnen einen modernen, attraktiven und stets aktuellen Zugang zu allen statistischen Informationen. Gerne weisen wir Sie auf folgende, besonders häufig genutzte Angebote hin.

Publikationsdatenbank – Publikationen zur vertieften Information

Fast alle vom BFS publizierten Dokumente werden auf dem Portal gratis in elektronischer Form zur Verfügung gestellt. Gedruckte Publikationen können bestellt werden unter der Telefonnummer +41 58 463 60 60 oder per Mail an order@bfs.admin.ch.
www.statistik.ch → Statistiken finden → Kataloge und Datenbanken → Publikationen

NewsMail – Immer auf dem neusten Stand



Thematisch differenzierte E-Mail-Abonnements mit Hinweisen und Informationen zu aktuellen Ergebnissen und Aktivitäten.
www.news-stat.admin.ch

STAT-TAB – Die interaktive Statistikdatenbank



Die interaktive Statistikdatenbank bietet einen einfachen und zugleich individuell anpassbaren Zugang zu den statistischen Ergebnissen mit Downloadmöglichkeit in verschiedenen Formaten.
www.stattab.bfs.admin.ch

Statatlas Schweiz – Regionaldatenbank und interaktive Karten



Mit über 4500 interaktiven thematischen Karten bietet Ihnen der Statistische Atlas der Schweiz einen modernen und permanent verfügbaren Überblick zu spannenden regionalen Fragestellungen aus allen Themenbereichen der Statistik.
www.statatlas-schweiz.admin.ch

Individuelle Auskünfte

Zentrale Statistik Information

+41 58 463 60 11, info@bfs.admin.ch

Um die Analysemöglichkeiten der Statistik der Unternehmensgruppen (STAGRE) zu erweitern, werden seit dem Erscheinungsjahr 2020 Resultate für den Umsatz bereitgestellt. Diese sind aggregiert nach verschiedenen Dimensionen (wie Grössenklasse, Art der Gruppe, Branche, Sitzland etc.) verfügbar. Dazu war es erforderlich, mehrere Quellen zu verknüpfen und, darauf aufbauend, statistische Modellierungen vorzunehmen.

Der vorliegende Bericht erläutert die Ausgangslage dazu sowie die Kriterien und Entscheidungen, die definiert und getroffen werden mussten, um die Variable «Umsatz» in der geforderten Qualität in der STAGRE verfügbar zu machen. Er beinhaltet sowohl ein Beschrieb des schlussendlich verwendeten Modells, als auch die technischen Details der Evaluation der verschiedenen in Betracht gezogenen Modellvarianten.

Online

www.statistik.ch

Print

www.statistik.ch
Bundesamt für Statistik
CH-2010 Neuchâtel
order@bfs.admin.ch
Tel. +41 58 463 60 60

BFS-Nummer

2239-2300

ISBN

978-3-303-06342-2

Die Informationen in dieser Publikation tragen zur Messung der Ziele für nachhaltige Entwicklung (SDG) bei.



Indikatorensystem MONET 2030

www.statistik.ch → Statistiken finden → Nachhaltige Entwicklung → Das MONET 2030-Indikatorensystem

**Statistik
zählt für Sie.**

www.statistik-zaehlt.ch