



Octobre 2023

Calcul des rabais des modèles – Possibilités d'évolution relatives à la circulaire 5.3

Contexte

Modèles : importance et ancrage légal

Également appelées « modèles », les formes particulières d'assurance impliquant un choix limité des fournisseurs de prestations¹ (p. ex. pour le premier point de contact) revêtent une grande importance. Près de trois quarts des assurés ont opté pour ce type de couverture d'assurance. Les modèles ont comme points communs d'améliorer l'efficacité des soins et d'engendrer moins de coûts qu'avec le libre choix du fournisseur de prestations (assurance de base) en cas de traitement équivalent. Partant, les personnes assurées selon les modèles payent une prime moins élevée².

L'OFSP est chargé d'évaluer les primes de l'AOS en Suisse³. Il exerce son activité de surveillance selon la circulaire [5.3](#), en relation avec la circulaire [5.1](#) (ch. 3.1.4). Ce document décrit la méthodologie appliquée pour calculer les rabais maximaux autorisés pour chaque modèle d'un assureur.

La circulaire 5.3 n'a pas été modifiée au cours des dernières années. Une modification pourrait devenir nécessaire suite à deux évolutions. D'une part, de plus en plus de personnes optent pour un modèle et, d'autre part, la compensation des risques avec les PCG a été affinée en 2020⁴. C'est pourquoi la question d'éventuellement adapter la circulaire 5.3 ponctuellement s'est posée. L'OFSP a chargé l'entreprise PricewaterhouseCoopers SA (PwC) d'élaborer un rapport en répondant à deux questions. Les points principaux sont résumés ci-après. Tout d'abord, le présent document décrit la méthodologie de base utilisée pour calculer les rabais d'un modèle.

¹ Cf. [art. 41, al. 4, LAMal](#), et [art. 99 OAMal](#)

² Cf. [art. 62, al. 1, LAMal](#)

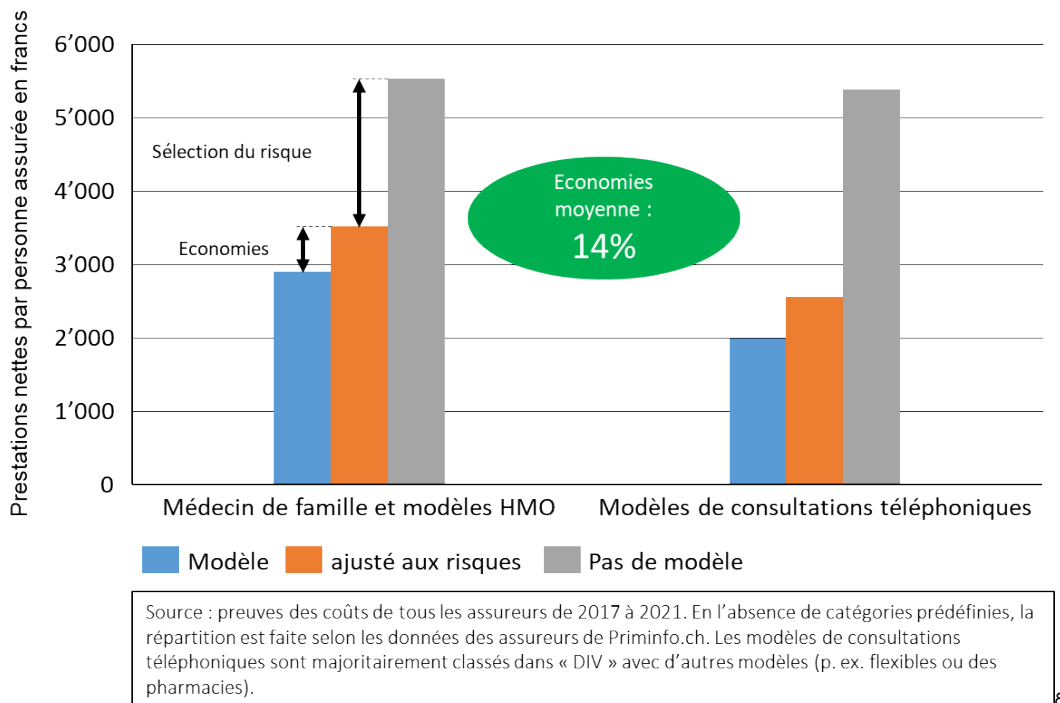
³ [Art. 16 LSAMal](#), et [art. 25 et 27 OSAMal](#)

⁴ PCG signifie « groupes de coûts pharmaceutiques » (*pharmaceutical cost groups*)
Cf. [site Internet de l'OFSP](#)

L'assurance de base, un groupe de référence pour calculer les rabais de modèle

Le Conseil fédéral fixe les limites maximales des réductions de primes en se fondant notamment sur les besoins de l'assurance⁵. Ces réductions ne sont admises que pour les différences de coûts qui résultent du choix limité des fournisseurs de prestations et non des différences dans les structures de risque⁶.

À cette fin, la circulaire 5.3 procède à une comparaison entre les personnes assurées selon l'assurance de base et celles ayant opté pour un modèle. Elle répartit ces deux groupes en classes égales (p. ex. âge, séjour à l'hôpital au cours de l'année précédente). L'on calcule ensuite, pour les personnes assurées selon un modèle, les prestations qui seraient dues si elles étaient assurées selon l'assurance de base et ce, pour chaque classe. Ces prestations hypothétiques sont ensuite comparées aux prestations effectives⁷. Cette preuve des coûts garantit que le montant du rabais du modèle ne dépend pas du fait que le modèle concerne majoritairement des personnes plus jeunes ou en meilleure santé.



Conformément à la circulaire 5.3, le rabais maximal autorisé est fixé de manière uniforme pour chaque modèle. Les dispositions prévues dans les circulaires 5.1 et 5.3 s'appliquent à tous les modèles. Les valeurs empiriques sont aussi indirectement prises en compte dans les nouveaux modèles (comparaison avec un modèle existant et, ainsi, appui indirect sur ses données empiriques au cas où le pourcentage demandé par l'assureur dépasse 14 %).

⁵ Cf. [art. 62, al. 3, LAMal](#)

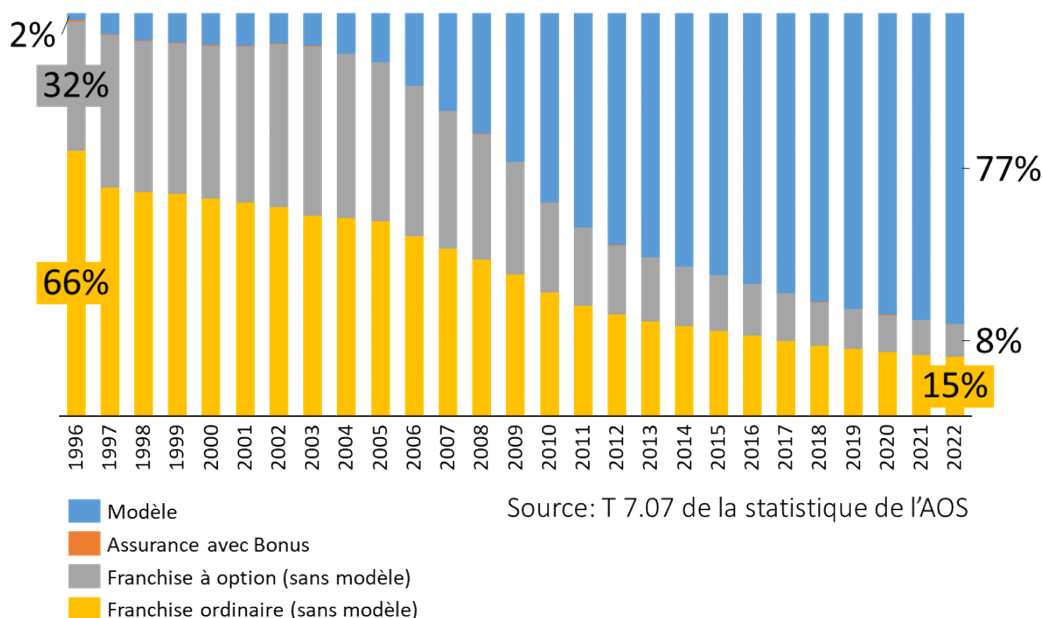
⁶ Cf. [art. 101 c, OAMal](#)

¹ La circulaire 5.3 se caractérise par le fait que les assureurs doivent respecter le rabais maximal calculé par l'OFSP en pourcents « uniquement » en moyenne (pondéré avec les effectifs) au niveau suisse. Il n'est pas autorisé de différencier le rabais du modèle des adultes. De même, chaque prime ne peut être inférieure à la prime minimale (règle des 50 %) visée à l'[art. 90c OAMal](#).

⁸ Cf. publication « Économies et sélection des risques dans les modèles d'assurance-maladie » 3/2013 de Jürg Burri, [lien](#)

Importance décroissante l'assurance de base

À l'heure actuelle, près de 25 % des personnes assurées peuvent choisir librement leur fournisseur de prestations. La proportion est en baisse depuis longtemps. Cette tendance devrait se poursuivre, probablement de façon moins marquée, étant donné que certaines de ces personnes ne veulent pas changer en raison du choix limité du fournisseur de prestations.



Méthode de calcul du rabais maximal : analyse par l'OFSP et dans le cadre d'une étude

Mandat d'étude

La preuve des coûts se fonde sur les données de cinq années. Les prestations nettes sont prises en compte par modèle et non par fournisseur de prestations, ce qui renforce la robustesse des calculs des rabais.

La section Primes et surveillance de la solvabilité a attribué un mandat concernant le calcul des rabais de modèle et le développement ou les alternatives relatifs à la circulaire 5.3 et mené un projet du 1^{er} octobre 2022 au 5 mai 2023. Avec l'intention de conserver la conception actuelle, l'étude vise à développer deux solutions alternatives qui portent sur les deux défis suivants posés dans la circulaire 5.3.

- **Défi I** : les effets du nouveau régime de la compensation des risques (introduction d'un critère PCG pour la prise en compte des prestations ambulatoires dans le cadre de l'harmonisation des structures de la preuve des coûts)
- **Défi II** : les effets de l'importance décroissante des personnes assurées ayant le libre choix du fournisseur de prestation (= groupe de référence pour les assurés avec une solution à modèle)

Solutions proposées dans l'étude

L'entreprise PwC a consigné ses constatations dans un rapport. Elle propose les deux solutions suivantes pour le **défi II** (réduction de l'importance de l'assurance de base) :

- a) le développement de la méthodologie actuelle figurant dans la circulaire 5.3 en s'appuyant sur les approches estimatrices de la « *Pooled Variance* » et de la variance totale⁹, au lieu de l'estimateur de variance actuel dans la circulaire 5.3 ;
- b) l'application d'une approche dite d'imputation fondée sur la régression pour enrichir les données concernant l'assurance de base

L'étude ne portait pas en priorité sur le **défi I** (PCG en tant que critères de classes), car la solution apportée au défi II permet de répondre également au défi I.

Le mandataire a évalué ses propositions à l'aune de l'approche existante visée dans la circulaire 5.3. À cette fin, il a appliqué les critères suivants :

- Nombre de classes applicables pour le calcul¹⁰ par rapport à toutes les classes disponibles
- Montant du rabais de modèle maximal autorisé
- Diverses mesures de variation en lien avec le rabais de modèle maximal autorisé

Première observation de l'étude

L'étude évalue les solutions a) et b) séparément puis de manière combinée. Le mandataire parvient toujours à la conclusion que cette approche permet de traiter le défi II (importance décroissante de l'assurance de base) et que l'approche actuelle du calcul des rabais maximums continue fondamentalement à produire des résultats valables, moyennant la mise en œuvre d'améliorations ciblées. Ce constat s'applique également si la tendance des effectifs en baisse dans l'assurance de base se poursuit¹¹.

Conformément aux critères d'évaluation ci-dessus, il recommande de combiner l'approche « *Total Variance* » et l'approche d'imputation via une régression linéaire. En conséquence, la robustesse accrue des résultats des modèles simplifie la solution pour faire face au défi I.

Les résultats de l'étude sur cette approche favorisée par le mandataire figurent à la page 41 du rapport, dans le tableau intitulée « (8) Total élargi I1 » et à l'annexe D.

Deuxième observation de l'étude

Dans le cadre de ses travaux, le mandataire s'est penché de manière approfondie sur le calcul actuel visé dans la circulaire 5.3.

En cas de mise en œuvre de l'approche concernant le défi II, l'étude prédit que le nombre minimum de deux assurés par classe, tel que la circulaire 5.3 le prescrit, peut être abaissé à un assuré. Ainsi, un plus grand nombre de classes peut être pris en compte pour le calcul du rabais de modèle maximal autorisé, ce qui permettrait d'augmenter la robustesse.

⁹ Estimateur dérivé d'ANOVA (= « Analysis of Variance »).

¹⁰ Cf. sous-chapitre plus haut « L'assurance de base, un groupe de référence pour calculer les rabais de modèle ». De manière simplifiée, une classe constitue par exemple la ligne Excel Année 2022, canton de Zurich, région 1, classe d'âge 41-45 ans, femme, franchise élevée, pas de séjour dans un hôpital ou un EMS l'année précédente, pas décédée en cours de l'année d'analyse.

¹¹ Cela ne s'applique pas si, d'un point de vue statistique, l'assurance de base n'a plus d'importance.

L'OFSP a notifié certaines incertitudes concernant le nombre d'assurés dans la circulaire 5.3. Le « Nombre d'assurés » n'est pas bien défini pour les personnes assurées pour moins de 12 mois (p. ex. arrivée dans un canton en cours d'année, départ d'un canton, décès ou naissance), mais il l'est pour celles assurées toute l'année. La définition de « Nombre d'assurés » ne faisant pas partie du mandat de projet, l'OFSP se penchera sur cette question ultérieurement.

Prochaines étapes de l'OFSP

Le rapport montre que l'idée maîtresse sous-tendant la circulaire 5.3, c.-à-d. comparer les personnes assurées selon un modèle avec celles ayant contracté une assurance de base, reste justifiée et produit des résultats suffisamment robustes.

Il apporte des propositions concrètes pour améliorer la stabilité des méthodes de calcul. À cette fin, il convient de procéder à des adaptations méthodologiques gérables. L'OFSP estime important de préserver dans une large mesure la stabilité des rabais de modèle maximaux autorisés et d'assurer en outre que les résultats des calculs fluctuent le moins possible chaque année.

Dans le cadre du mandat, des calculs utilisant le langage R ont été effectués et testés. Le code R développé par l'OFSP a été complété par les solutions a) et b) mentionnées ci-dessus pour le défi II (importance décroissante de l'assurance de base).

L'OFSP est convaincu que les constats du mandataire portant sur la circulaire 5.3 ainsi que ses propositions de développement peuvent engendrer des améliorations fondamentales. Le choix et la mise en œuvre des approches de solutions sont en cours d'évaluation. Dans cette optique, l'OFSP collaborera étroitement avec la branche pour disposer à l'avenir d'un ensemble de règles compréhensible et dont le contenu est fondamentalement acceptable.

Bundesamt für Gesundheit BAG

Berechnung der
Modellrabatte –
Weiterentwicklung oder
Alternativen betreffend
Kreisschreiben 5.3

5. Mai 2023

Inhaltsverzeichnis

1. Zusammenfassung	3
1.1. Zweck, Vorgehen und Ergebnisse	3
1.2. Zusammenfassung der Erkenntnisse	4
2. Einleitung	6
2.1. Auftraggeber und Auftragsbeschreibung	6
2.2. Ausgangslage und Zielsetzung	6
2.3. Verantwortung des BAG	9
2.4. Autoren	9
2.5. Aufsichtsrechtliche Vorgaben und relevante Referenzen	10
2.6. Algorithmus und Codierung	10
3. Diskussion des bestehenden Ansatzes	12
3.1. Beschreibung des bestehenden Ansatzes (KS 5.3)	12
3.2. Validierung des bestehenden Ansatzes (KS 5.3)	16
4. Entwicklung alternativer Ansätze	21
4.1. Pooled- und Totale-Varianzschätzer	21
4.2. Imputationsansatz	26
5. Bewertung der Ansätze	32
5.1. Qualitative Bewertung	32
5.2. Quantitativer Ansatz zur Bewertung	34
5.2.1. Bootstrap-Verfahren	34
5.2.2. Definition der Gütekriterien	36
5.2.3. Datengrundlage	36
5.3. Quantitative Bewertung	39
5.3.1. Definition der bewerteten Ansätze/Ansatzkombinationen	39
5.3.1. Auswertungen	41
6. Analyse Pharmazeutische Kostengruppen (PCG)	44
6.1. Einführung der Kostengruppen	44
6.2. Bewertung der Ansätze bei Einbezug der Kostengruppen (PCG)	44
Appendix A. - Abkürzungen	46
Appendix B. - BAG-Daten und Lieferergebnisse	47
Appendix C. - Beweise	48
Appendix D. - Weitere Analyseergebnisse	52
Referenzen	54

1. Zusammenfassung

1.1. Zweck, Vorgehen und Ergebnisse

In diesem Bericht werden die Vorgehensweise und Ergebnisse zum Projekt «Berechnung der Modellrabatte – Weiterentwicklung oder Alternativen betreffend Kreisschreiben 5.3» beschrieben.

Die Sektion Prämien und Solvenzaufsicht (PuS) ist beim Bundesamt für Gesundheit (BAG) für die Beaufsichtigung der Modelle zuständig. Sie übt ihre aufsichtsrechtliche Tätigkeit via Kreisschreiben (KS) 5.3 i.V. mit 5.1 (Ziffer 3.1.4) aus. In Bezug zu Art. 101 KVV teilt das KS 5.3 die Versicherten der Basisversicherung und der Modelle in gleiche Klassen ein (z.B. Alter, Spitalaufenthalt im Vorjahr, Franchise). Dabei wird für die Modellversicherten jeder Klasse berechnet, welche Leistungen anfielen, wenn sie in der Basisversicherung versichert wären und mit ihren effektiven Leistungen verglichen. Das BAG hat im KS 5.3 zwei Schwachstellen identifiziert, welche in den folgenden Bereichen liegen:

- **Schwachstelle (1):** Der Kostennachweis (KN) gemäss KS 5.3 stützt sich nicht auf die neueste Entwicklung im Bereich Risikoausgleich ab. Im KN gibt es derzeit kein Kriterium zu ambulanten Leistungen/Konsum von Medikamenten (resp. Wirkstoffen). Das BAG vermutet, dass dadurch die Bestandesstrukturangleichung zwischen Basisversicherten (freie Wahl des Leistungserbringers) und Modellversicherten beeinträchtigt ist. Dies wiederum würde die Aussagekraft der maximal zulässigen Modellrabatte mindern.
- **Schwachstelle (2):** In der Berechnung des maximal zulässigen Modellrabatts gemäss KS 5.3 werden die Basisversicherten als Vergleichskollektiv herangezogen. Die Historie zeigt, dass die Anzahl der Basisversicherten zu Gunsten der Modellversicherten jedoch seit Jahren abnimmt und ein Ende dieses Trends nicht abzusehen ist. Dies führt dazu, dass es zu einem Ungleichgewicht zwischen Modellversicherungskollektiv und Vergleichskollektiv kommt. Dies macht die Berechnungen des BAG eventuell instabil und die Aussagekraft (v.a. bei Versicherern mit hohem Anteil an Modellversicherten) des Ansatzes wird eingeschränkt.

Das BAG sucht nach zwei alternativen Lösungswegen, welche die beiden Schwachstellen im KS 5.3 ansprechen und somit künftig die Festlegung der Modellrabatte stabil vorgenommen werden kann.

Da die unmittelbare Berücksichtigung von Pharmazeutischen Kostengruppen (PCG) im Kostennachweis gemäss KS 5.3 die beschriebene Schwachstelle (2) zunächst verstärkt¹, wird die Behebung der Schwachstelle (2) als Grundvoraussetzung zur Lösung der Schwachstelle (1) gesehen. Daher werden zwei alternative Ansätze erarbeitet, welche vor allem die Schwachstelle (2) adressieren und es werden dabei die bisherige Konzeption im gegenwärtigen KS 5.3 zur Bestimmung des maximal zulässigen Modellrabatts mittels eines Vergleichs von Kollektiven weitestgehend erhalten:

- a) Weiterentwicklung der bestehenden Methodik im Kreisschreiben 5.3 mit alternativen Schätzern der Art «Pooled Varianz» und «Totale Varianz» Schätzer, welche die aktuellen Varianzschätzer im KS 5.3 ersetzen; und
- b) Anwendung eines Imputationsansatzes basierend auf Regression zur Anreicherung von Datenpunkten in schrumpfenden Beständen.

Um die Vorteile der beiden Ansätze a) und b) berücksichtigen zu können, werden zudem verschiedene Kombinationen als alternative Lösungsansätze gebildet (z.B. Totale Varianz mit Imputation). Die erarbeiteten alternativen Lösungsansätze werden mit verschiedenen qualitativen und quantitativen Beurteilungskriterien gegenüber dem bestehenden Schätzverfahren im KS 5.3 zur Messung der zulässigen

¹ Da eine oder mehrere zusätzliche Kategorisierungsvariable(n) für die «PCG» in den Kostennachweis (KN) gemäss KS 5.3 eingefügt werden müssen und somit weniger Versicherte auf die definierten Klassen zugeordnet werden.

Modelrabatte ausgewertet. Das bestehende Schätzverfahren im KS 5.3, wie es aktuell in Anwendung ist, definieren wir als Referenzmodell. Für die Beurteilung und einem Vergleich der Leistungsfähigkeit der Ansätze ist die Verwendung von objektiven Gütekriterien essenziell. Neben deskriptiven Beurteilungskriterien wird die Schätzunsicherheit des Maximalrabatts (R_{max}) unter den verschiedenen Ansätzen quantitativ beurteilt. Um die Ungenauigkeit in der Schätzung evaluieren zu können, wurde die Varianz für die Schätzung $\hat{R}_{max}$ (Statistik) als Streuungsmass gewählt. Ein weiteres zentrales Kriterium war der Erhalt der Stabilität (Veränderung) von $\hat{R}_{max}$ gegenüber jener bestimmt mit dem Referenzmodell. Diese Analysen werden durch einen sogenannten Bootstrap-Ansatz auf Basis von gewählten Testdatensätzen ausgeführt.

Gemäss unseren Analysen und erhaltenen Ergebnissen zur Messung der Leistungsfähigkeit der Schätzer zeigt sich, dass sämtliche alternative Ansätze die Schwachstelle (2) adressieren können. Dies führt dazu, dass die Menge der verwendeten Datensätze in den alternativen Ansätzen im Vergleich zum Referenzmodell (teils signifikant) erhöht werden kann. Eine Kombination aus «Totaler Varianz mit Imputation» scheint am vielversprechendsten zu sein, um die Schwachstelle (1) und (2) zu adressieren. Es zeigt sich, dass die Einführung der PCG als zusätzliches Klassenkriterium die zur Verwendung stehenden Daten im bestehenden Referenzmodell nach KS 5.3 weiter merklich einschränkt. Der bestehende Ansatz im KS 5.3 stellt also mittel- bis langfristig keine geeignete Option dar, wenn PCG explizit in Form von zusätzlichen Kategorien berücksichtigt werden sollen. Einen eindeutigen Unterschied bezüglich der Art der PCG-Segmentierung scheint es gemäss unseren Analysen nicht zu geben. Es fällt jedoch auf, dass die Höhe des Maximalrabatts R_{max} bei der Einführung von PCG sich deutlich verändert. Die Höhe der Veränderungen scheinen allerdings unabhängig vom verwendeten Ansatz zu sein. Es ist allerdings hervorzuheben, dass bei Verwendung des alternativen Ansatzes «Totale Varianz mit Imputation» wieder mehr Klassen zur Verfügung stehen, was eine Grundvoraussetzung für den Einbezug von PCG in die Berechnung des Kostennachweises ist.

Die vorgeschlagenen Ansätze bieten Lösungen, welche die erforderlichen Auswertungen gemäss KS 5.3 bei kleineren Bestandsgrössen stabilisieren können. Unter den analysierten Lösungsvorschlägen hat sich der Ansatz «Totaler Varianz mit Imputation» als stabil und praktikabel in der Umsetzung erwiesen. Die Aussagekraft der Auswertungen wird erhalten, auch wenn sich der Trend von abnehmenden Basisbeständen in den nächsten Jahren erwartungsgemäss weiter fortsetzen wird. Die vorgestellten Ansätze stellen jedoch auch dann nur bedingte Lösungen für den Fall dar, dass die Basisversicherung und somit das notwendige Vergleichskollektiv zu einem späteren Zeitpunkt aus statistischer Perspektive keinerlei Bedeutung mehr darstellen sollte. In einem derartigen Szenario wird empfohlen, die grundsätzliche Konzeption des KS 5.3 rechtzeitig zu überarbeiten bzw. gar die Definition des Vergleichskollektivs anzupassen.

Unsere Haupteckentnis und weitere Beobachtungen werden in folgender Sektion zusammengefasst.

1.2. Zusammenfassung der Erkenntnisse

Zur Bearbeitung der identifizierten Schwachstellen im gegenwärtigen KS 5.3, vor allem der Schwachstelle (2), stellen wir die folgenden Resultate vor.

Haupteckentnis – Verwendung des alternativen Ansatzes «Totale Varianz mit Imputation»

Gemäss unseren Analysen und erhaltenen Ergebnissen zur Messung der Leistungsfähigkeit der Schätzer zeigt sich, dass sämtliche alternativen Ansätze die Schwachstelle (2) adressieren können. Das führt dazu, dass die Menge der verwendeten Datensätze in den alternativen Ansätzen im Vergleich zum Referenzmodell (teils signifikant) erhöht werden kann. Allein auf Grund der theoretischen Basis ist die Implementierung eines Ansatzes, welcher auf der Totalen Varianz basiert, sinnvoll. Dieser Ansatz erfordert die wenigsten Restriktionen an die zugrundeliegenden Daten. Wird der Ansatz nach Totaler Varianz noch mit einer

Imputation kombiniert, sollte auch eine adäquate Lösung für die Schwachstelle (1) ermöglicht werden. Aus unserer Sicht wäre hier der Ansatz I1 «Log-Lineare Regression ohne vorgelagerte Logistische Regression» zu präferieren, da dieser – gegeben der Datenbasis – deutlich stabilere Resultate generiert als auch leichter gegenüber der Versicherungswirtschaft zu kommunizieren sein dürfte als der Imputationsansatz I2 («Log-Lineare Regression mit vorgelagerter Logistischer Regression»).

Beobachtung #1 – Präzision der Beschreibungen (Notation) im KS 5.3

Die Beschreibungen im KS 5.3 sind an einigen Stellen unpräzise. Beispielsweise wird für die «Anzahl Versicherte» (d.h., NMC_k und $NBase_k$) im gegenwärtigen KS 5.3 eine rationale Zahl zugelassen (siehe auch Tabelle 3). In den Formeln auf S. 7 des KS 5.3 wird dieser Wert als Summenobergrenze in den Berechnungen für die Variable LMC_k verwendet. Dies ist weder mathematisch wohldefiniert (Summationen benötigen ganzzahlige Grenzen), noch kann es in verschiedenen Fällen praktisch ausgewertet werden. Hinzu kommt, dass es nicht der Definition der LMC_k auf S. 6 des KS 5.3 entspricht. Würden diese Ungenauigkeiten durch eine einheitliche und mathematisch widerspruchsfreie Definition behoben werden, dann würde der angepasste Schätzer $Var(A)$ von der im KS 5.3 (S. 7 als Verweis auf die Definition) definierten Formel abweichen. Entsprechende Anpassungen wären ebenfalls für die Formeln von B und $Var(B)$ nötig.

Neben den mathematisch unpräzisen Ableitungen der Varianz-Formeln, könnte die nicht wohldefinierte Darstellung im KS 5.3 (zwischen Definitionen und Formeln) auch dazu führen, dass Versicherungsunternehmen die Beschreibung «Anzahl Versicherte» missverstehen, vom naheliegenden ganzzahligen Wert ausgehen und somit möglicherweise ungeeignete Informationen im EF-MC Datensatz erfassen und an das BAG übermitteln.

Aus den oben genannten Ungenauigkeiten ergaben sich im Folgenden immer wieder Herausforderungen bei der Definition alternativer Ansätze als auch bei den Auswertungen der entsprechenden Tests. In den folgenden Abschnitten werden wir mit den im KS 5.3 definierten/formulierten Schätzern weiterarbeiten, als wenn diese wohldefiniert wären, um unter anderem alternative Ansätze mit dem aktuellen Vorgehen im KS 5.3 vergleichen zu können. Entsprechend unterliegen die vorgestellten Alternativansätze den gleichen Limitationen wie dies im aktuellen KS 5.3 der Fall ist.

Beobachtung #2 – Definition der Schätzer A und B im KS 5.3 als (gewichtete) arithmetische Mittelwerte

Die Definition der Durchschnittsleistungen A und B beruht im gegenwärtigen KS 5.3 nicht auf arithmetischen Mitteln. Dies hat zur Folge, dass die resultierenden Varianzen der Schätzer A und B nicht unmittelbar auf Basis von statistischen Definitionen in der Fachliteratur abgeleitet werden können. Entsprechende Eigenschaften von Mittelwerten und Varianzen sind somit nicht immer allgemein gültig, auch wenn diese im KS 5.3 generell zur Anwendung kommen.

2. Einleitung

2.1. Auftraggeber und Auftragsbeschreibung

Der Auftrag wurde von der Schweizerischen Eidgenossenschaft, Bundesamt für Gesundheit BAG, vertreten durch den Direktionsbereich Kranken- und Unfallversicherung, Abteilung Versicherungsaufsicht, Sektion Prämien und Solvenzaufsicht (PuS) erteilt. Der Auftrag mit dem Titel «Berechnung der Modellrabatte - Weiterentwicklung oder Alternativen betreffend dem Kreisschreiben 5.3» wurde am 29. September 2022 mit dem Zweck erteilt, die bestehende Vorgehensweise zur Berechnung der Modellrabatte wie im gegenwärtigen Kreisschreiben 5.3 dargestellt, zu analysieren und zwei alternative Lösungswege zur Adressierung bekannter Schwachstellen im Kreisschreiben 5.3 auszuarbeiten. Hierzu wurden grundsätzlich die folgenden beiden Arten von Ansätzen ausgewählt:

- a) Weiterentwicklung der bestehenden Methodik im Kreisschreiben 5.3 mit alternativen robusteren Schätzern, und
- b) Verwendung einer Multiple Linearen Regression bzw. eines Imputationsansatzes.

Die Möglichkeit der Aggregation von Klassen oder deren Umsegmentierung wurde in Betracht gezogen, jedoch auf Grund erwarteter erheblicher Änderungen in den Ergebnissen zu den zulässigen Maximalrabatten gemäss KS 5.3 nicht weiterverfolgt (siehe Bemerkung 1, Punkt (c) (ii) auf Seite 9). Um die Stabilität der beiden alternativen Ansätze bei geringen Beständen (v.a. in der Basisversicherung) zu erhöhen, wurde auf Techniken der Datenimputation («Anreicherung von Datenpunkten») zurückgegriffen.

Die Implementation der weiterentwickelten oder alternativen Ansätze zur Berechnung der Modellrabatte gemäss dem KS 5.3 wurde in der Programmiersprache **R** vorgenommen. Als Ausgangsbasis stellte das BAG den bestehenden R-Code zum derzeitigen Kostennachweis gemäss KS 5.3 zur Verfügung. Dazu gehörte auch ein veränderter und somit nicht-repräsentativer Testdatensatz. Damit konnte PwC Schweiz (nachfolgend PwC) die für den Auftrag relevanten Teile des R-Codes ausführen.

2.2. Ausgangslage und Zielsetzung

Das BAG (Sektion Prämien und Solvenzaufsicht) ist gemäss Art. 16 KVAG i.V. mit Art. 7 KVAG und Art. 25 und Art. 27 KVAV zuständig für die Prüfung der Prämien der OKP Schweiz. Besonderen Versicherungsformen mit eingeschränkter Wahl der Leistungserbringer gemäss Art. 41 Abs. 4 KVG und Art. 99 KVV (sog. Modelle) kommt eine grosse Bedeutung zu (ca. drei Viertel der Versicherten sind momentan in einem Modell versichert). Für die Versicherer sind Modelle wichtig, weil sie sich damit nicht nur preislich, sondern in beschränktem Ausmass auch inhaltlich von der Konkurrenz abheben können. Volkswirtschaftlich sind die Modelle zentral, weil sie die Versicherten mit Leistungen zu kostensparenderem Verhalten veranlassen. Diese profitieren im Gegenzug von einer tieferen Prämie (Art. 62 Abs. 1 KVG).

Gemäss Art. 62 Abs. 3 KVG legt der Bundesrat insbesondere aufgrund versicherungsmässiger Erfordernisse Höchstgrenzen für die Prämienermässigung fest. In Art. 101 KVV ist verankert, dass Prämienermässigungen nur zulässig für Kostenunterschiede sind, die auf die eingeschränkte Wahl der Leistungserbringer sowie auf eine besondere Art und Höhe der Entschädigung der Leistungserbringer zurückzuführen sind.

Die Sektion Prämien und Solvenzaufsicht (PuS) ist im BAG für die Beaufsichtigung der Modelle zuständig. Sie übt ihre aufsichtsrechtliche Tätigkeit via Kreisschreiben (KS) 5.3 i.V. mit 5.1 (Ziffer 3.1.4) aus. Das KS teilt mit Blick auf Art. 101 KVV die Versicherten der Basisversicherung und der Modelle in gleiche Klassen ein (z.B. Alter, Spitalaufenthalt im Vorjahr, Franchise). Es wird für die Modellversicherten jeder Klasse

berechnet, welche Leistungen anfielen, wenn sie in der Basisversicherung versichert wären und mit ihren effektiven Leistungen verglichen.²

Dem BAG steht ergänzend als Regulierungsinstrument Art. 101 KVV zur Verfügung. Dieser hält fest, dass die Versicherten der Modelle keine besonderen Risikogemeinschaften sind. Der Versicherer hat bei der Prämienfestsetzung darauf zu achten, dass die Modelle im versicherungstechnisch erforderlichen Mass an die Reserven beitragen. Das BAG ist somit befugt, die finanzielle Situation der Modelle ex-post zu analysieren (z.B. via Kennzahl Combined Ratio).

Schwachstellen im Aktuellen KS 5.3:

Das KS 5.3 ist zwar von der Branche grundsätzlich gut akzeptiert, dennoch enthält es zwei Punkte, für die sich das BAG im Rahmen dieses Auftrages eine detailliertere Analyse bzw. die Definition von Handlungsempfehlungen gewünscht hat. Das KS 5.3 soll in den folgenden zwei Bereichen weiterentwickelt werden:

- **Schwachstelle (1):** Der Kostennachweis (KN) gemäss KS 5.3 stützt sich nicht auf die neuste Entwicklung im Bereich Risikoausgleich ab. Dieser wurde 2020 stark überarbeitet. Der sehr einfach aufgebaute Medi-flag wurde durch sog. PCG (vgl. als Informationsquelle: Homepage BAG, Gemeinsame Einrichtung KVG) abgelöst³. Im KN gibt es derzeit kein Kriterium zu ambulanten Leistungen/Konsum von Medikamenten (resp. Wirkstoffen). Das BAG vermutet, dass dadurch die Bestandesstrukturangleichung zwischen Basisversicherten (freie Wahl des Leistungserbringers) und Modellversicherten beeinträchtigt ist. Dies wiederum würde die Aussagekraft der maximal zulässigen Modellrabatte mindern.
- **Schwachstelle (2):** In der Berechnung des maximale zulässigen Modellrabatts gemäss KS 5.3 werden die Basisversicherten als Vergleichskollektiv herangezogen. Die Historie zeigt, dass die Anzahl der Basisversicherten zu Gunsten der Modellversicherten jedoch seit Jahren abnimmt und ein Ende dieses Trends nicht abzusehen ist. Dies führt dazu, dass es zu einem Ungleichgewicht zwischen Modellversicherungskollektiv und Vergleichskollektiv kommt, was die Berechnungen des BAG eventuell instabil macht und so die Aussagekraft (v.a. bei Versicherern mit hohem Anteil an Modellversicherten) des Ansatzes einschränkt.

Das BAG sucht nach Lösungen, wie die Festlegung der Modellrabatte robuster gemacht werden könnte (unter Beachtung von Art. 62 Abs. 3 KVG und bei neuen Modellen von Art. 101 KVV). Dazu gehören auch die Modellierung und Anwendung der Streuung in der Formel zur Berechnung des maximalen Modellrabatts.⁴

Bemerkung 1:

- (a) In den folgenden Abschnitten werden alternative Ansätze diskutiert, welche vor allem die Schwachstelle (2) adressieren. Diese Ansätze beziehen sich vornehmlich auf das aktuelle Konzept im KS 5.3 und zielen darauf ab, nur geringfügig in die bestehende Herangehensweise eingreifen zu müssen (z.B. durch Ersetzen der Definitionen für Varianzen). Die vorgeschlagenen Ansätze bieten

² Eine Spezialität des KS 5.3 ist, dass die Krankenversicherer den vom BAG berechneten maximalen Rabattsatz in Prozenten «nur» im Durchschnitt (mit den Beständen gewichtet) auf Stufe CH einhalten müssen. Nicht erlaubt ist hingegen, den Modellrabatt der Erwachsenen zu differenzieren. Auch darf die einzelne Prämie keinesfalls die sog. Minimale Prämie (50%-Regel) nach Art. 90c KVV unterschreiten.

Exkurs: Das KS 5.3 verlangt die Einteilung der Modelle in solche ohne Vertrag zwischen Versicherer und Leistungserbringer und bei HMO (im Kostennachweis KS 5.3 sind bei HMO auch Hausarztmodelle enthalten) und DIV in solche mit Vertrag ohne/mit Budgetmitverantwortung. Auf priminfo (Prämien-Vergleichsplattform des BAG) werden Hausarztmodelle separat von HMO-Modellen betrachtet, was immer fraglicher erscheint (vgl. Tätigkeitsbericht, S. 21). Die Verträge zwischen Krankenversicherern und Leistungserbringern sind nicht öffentlich zugänglich. Für Dritte einsehbar sind hingegen die Vertragsbedingungen üblicherweise je Modell, welche die Versicherer i.d.R. auf ihrer Homepage platzieren.

³ Es wird angemerkt, dass der Medi-flag nur eine Übergangslösung im Risikoausgleich ohne PCG bis heute mit PCG darstellt.

⁴ Anwendung: Auseinandersetzung, wie die Varianz in die Berechnung des Maximalrabatts einfließt (vgl. S. 9 des KS 5.3)
Modellierung: Definition und (statistische) Schätzung der Varianz (vgl. S. 7 des KS 5.3)

zudem Lösungen, die erforderlichen Auswertungen gemäss dem KS 5.3 bei kleineren Bestandsgrössen zu stabilisieren und somit deren Aussagekraft zu erhalten, obwohl sich der Trend von abnehmenden Basisbeständen in den nächsten Jahren erwartungsgemäss weiter fortsetzen wird. Diese Ansätze sind nur bedingt mögliche Lösungen für den Fall, dass die Basisversicherung und somit das notwendige Vergleichskollektiv zu einem späteren Zeitpunkt aus statistischer Perspektive keinerlei Bedeutung mehr haben sollte. In einem derartigen Szenario wird empfohlen, die grundsätzliche Idee des KS 5.3 zu überarbeiten oder gar die Definition des Vergleichskollektivs zu ersetzen.

- (b) Bei Schwachstelle (1) soll eine oder mehrere zusätzliche Kategorisierungsvariable(n) für die «PCG» in den Kostennachweis (KN) gemäss KS 5.3 eingefügt werden. Im aktuellen KS 5.3-Konzept würde die Aufnahme von weiteren Variablen für PCG dazu führen, dass die bereits stark schrumpfenden Basisbestände⁵ in den bestehenden Kategorien noch weiter aufgeteilt werden müssen. Damit würde die Problematik der Schwachstelle (2) unmittelbar verschärft werden. Entsprechend ist die Behebung der Schwachstelle (2) eine Grundvoraussetzung bevor Schwachstelle (1) im aktuellen Konzept des KS 5.3 sinnvoll angegangen werden kann. Der Vorteil an dem Vorgehen zuerst Schwachstelle (2) zu bearbeiten besteht darin, dass als Konsequenz auch Schwachstelle (1) adressiert werden kann, ohne von der grundsätzlichen Konzeption des KS 5.3 abweichen zu müssen. Eine grundsätzliche Anpassung der Methodik wäre also auch bei Einbezug der Kostengruppen (Schwachstelle 1) somit nicht zwingend nötig.
- (c) Weitere Möglichkeiten die Schwachstelle (2) zu bearbeiten und damit die Grundvoraussetzung für die Behebung von Schwachstelle (1) zu legen, kann durch folgende Ansätze erfolgen:
 - (i) *Multiple Linear Regression auf Versicherten-Ebene (Regression auf EFIND)*: Dieser Ansatz hätte die grösste Ähnlichkeit zur Vorgehensweise im Risikoausgleich. Er würde auch die Möglichkeit einer analogen Behandlung von Schwachstelle (1) wie im Risikoausgleich bieten und hätte darüber hinaus weitere Vorteile wie zum Beispiel: keine Abhängigkeit von Klassendefinitionen, geringere Dateneinschränkungen, Rabatt-Information liegen auf Versicherten-Ebene vor, und bei der Messung des Einflusses von Klassenkriterien auf die Höhe der Modellrabatte. Ein Ansatz auf Versichertenebene erfordert den Zugriff auf EFIND-Daten, welche jedoch erst zu einem späteren Zeitpunkt zur Verfügung stehen würden, als PuS diese für die Prüfungshandlungen zur Höhe der Modellrabatte benötigt. Zudem müsste das KS 5.3 – auch im grundsätzlichen Konzept – essentiell überarbeitet werden. Dieser Ansatz wurde mit den Vertretern des PuS im Herbst 2022 diskutiert. Da zudem das KS 5.3 gänzlich ersetzt werden müsste, wurde dieser weniger favorisiert und daher ein derartiger Ansatz nicht weiterverfolgt oder detailliert.
 - (ii) *Aggregation/Umsegmentierung*: Schwachstelle (2) kommt in höherem Masse zum Tragen, wenn im aktuellen Ansatz des KS 5.3 die Basisversicherten auf sehr viele verschiedene Klassen verteilt werden müssen. Eine Möglichkeit wäre, dass einzelne Klassenkriterien abgelöst (z.B. Einschränkung nur auf Merkmale, die auch im Risikoausgleich enthalten sind) oder Attribute aggregiert werden (z.B. anstelle von 5-Jahres Altersgruppierungen eine Einteilung nur in Kinder/Junge Erwachsene/Erwachsene ggfs. in Kombination mit Unfalldeckung). Beide Ansätze würden die Anzahl an Klassen im aktuellen KS 5.3-Ansatz signifikant reduzieren, sodass die Einzelbestände in den Klassen grösser würden und somit auch die Schwachstelle (2) adressiert wird. Testrechnungen im Herbst 2022 haben gezeigt, dass bei der Aggregation bzw. Umsegmentierung von Klassen es zu erheblichen Änderungen in den Ergebnissen zu den zulässigen Maximalrabatten gemäss KS 5.3 pro Modell/Krankenkasse kommt und allenfalls einen hohen

⁵ Dieses Verhalten wird auch für die Bestände der Alternativmodelle beobachtet.

Kommunikationsaufwand als auch ungewünschte Verwerfungen in den Ergebnissen zur Folge hätte. Als Beispiel seien die Ergebnisse in Kapitel 6 zu nennen. Es wird ersichtlich, dass auf Grund der Einführung der Kostengruppen PCG sich die Maximalrabatte deutlich im Vergleich zur bestehenden Segmentierung/Aggregation ohne Merkmale für PCG verändern können. Um eine Konsistenz in der Höhe der Maximalrabatte über die Zeit sicherstellen zu können, müsste eine künstliche Glättung eingefügt werden⁶. Entsprechend wurde dieser Ansatz im Rahmen dieses Mandats nicht weiterverfolgt oder detailliert.

- (d) Während der Projektarbeit an der Methode Varianz traten teilweise unplausible Ergebnisse auf. Die Ursache liegt im Referenzmodell und das BAG übernahm den Lead für die Aufklärung. Beim Referenzmodell kamen Fragen zur Definition Anzahl Versicherten auf. Für ganzjährig Versicherte ist die Anzahl Versicherter wohldefiniert, für kürzere Verweildauern hingegen nicht. Die Thematik der Definition Anzahl Versicherter war nicht Teil des Projektauftrages und wird vom BAG im Nachgang aufgegriffen.

2.3. Verantwortung des BAG

Die ordnungsgemässe Prüfung und damit die verbundenen Berechnungen der Modellrabatte als auch die Nutzung des R-Codes liegt im Zusammenhang mit Art. 16 KVAG i.V. mit Art. 7 KVAG und Art. 25 und Art. 27 KVAV beim BAG (Sektion Prämien und Solvenzaufsicht). Gemäss gesetzlichen Vorgaben ist das BAG für die Prüfung der Prämien der OKP Schweiz zuständig.

Für die Ausführung unserer Tätigkeiten im Rahmen der beschriebenen Vorgehensweisen und Ergebnisse in diesem Bericht, war das BAG verantwortlich für die:

- Bereitstellung eines lauffähigen R-Code zum derzeitigen Kostennachweis gemäss KS 5.3,
- Bereitstellung von veränderten und somit allenfalls nicht-repräsentativen Testdatensätzen, und die
- Bereitstellung von nicht repräsentativen EFIND-Testdatensätzen, so dass

wir die relevanten Teile des bestehenden R-Codes anpassen konnten. Um die Funktionalität des angepassten R-Codes und der Lösungsansätze nachhaltig testen zu können, hat das BAG weitere Testläufe auf echten Daten ausgeführt. Dies erfolgte in Verantwortung des BAG, während PwC über die Resultate der entwickelten Methoden und des R-Codes laufend informiert wurde, um weitere Verbesserungen und Korrekturen vornehmen zu können. Die PwC bzw. die nachfolgend aufgeführten Autoren und Projektmitglieder hatten selbst keinen Zugang zu echten Datensätzen des BAG.

2.4. Autoren

Die Tätigkeiten zur Entwicklung der alternativen Lösungsansätze, Implementierungen im R-Code und Bereitstellung der Ergebnisse wurde seitens PwC von folgenden Personen ausgeführt:

Nr.	Kontaktperson	Rolle	Qualifikation
1	Herr Dr. Harald Dornheim, Director	Projektleitung, Experte Statistik und Modellierung	Ph.D. Mathematische Statistik, Aktuar SAV, DAV, FRM, CERA
2	Herr Dr. Jan Wirfs, Manager	Stv. Projektleiter, Experte Statistik und Modellierung	Ph.D. HSG, Quantitative Versicherungswirtschaft, Aktuar SAV
3	Herr Dr. Richard Polifka, Senior Consultant	Projektmitglied, Implementierungen in R, Datenaufbereitung	Ph.D. Universität Prag, Teilchenphysik, Data Analytics

⁶ Die Glättung mittels Einbezugs von mehreren Jahren würde aus unserer Sicht keine zwingende Abhilfe darstellen.

4	Frau Alice Montoschi, Senior Consultant	Projektmitglied, Experte Krankenversicherungen	M.Sc. HEC Aktuarwissenschaften, Aktuarin SAV
---	--	---	--

Tabelle 1: Autoren und Projektmitglieder

Besprechungen fanden mit folgenden Vertretern des BAG aus dem Direktionsbereich Kranken- und Unfallversicherung, Abteilung Versicherungsaufsicht, Sektion Prämien und Solvenzaufsicht statt:

Nr.	Kontaktperson	Rolle
1	Herr Fabrice Perler	Umfassende Projektverantwortung, Teilnahme an gewählten Fachdiskussionen zur Konzeption der Methodiken und Besprechung Zwischenresultate
2	Frau Stefanie Zihler	Projektleitung BAG Modellrabatte, Führung Protokoll, Teilnahme an Fachdiskussionen zur Konzeption der Methodiken, Implementierungen in R und Besprechung Resultate
3	Frau Annina Stohrer-Nef	Stv. Projektleiterin, Teilnahme an Fachdiskussionen zur Konzeption der Methodiken, Implementierungen in R und Besprechung Resultate
4	Herr Oliver Stoecklin	Projektmitarbeiter, Teilnahme an Fachdiskussionen zur Konzeption der Methodiken, Implementierungen in R und Besprechung Resultate

Tabelle 2: Kontaktpersonen und Projektmitglieder

2.5. Aufsichtsrechtliche Vorgaben und relevante Referenzen

Für die Auftragserfüllung wurde unter anderem folgende aufsichtsrechtliche Vorgaben und Referenzen benützt:

- [KVG] Bundesgesetz über die Krankenversicherung (KVG) vom 18. März 1994 (Stand am 18. März 2023), insbesondere Art. 41 Abs. 4 KVG, Art. 62 Abs. 1 und 3 KVG.
- [KVV] Verordnung über die Krankenversicherung vom 27. Juni 1995 (Stand am 1. Januar 2023), insbesondere Art. 99 KVV, Art 101 KVV (sog. «Modelle»).
- [KS5.1] Kreisschreiben 5.1 «Prämien der obligatorischen Krankenpflegeversicherung und der freiwilligen Einzeltaggeldversicherung» (speziell Ziffer 3.1.4), Inkrafttreten 1. Juni 2022, vom 3. Mai 2022.
- [KS5.3] Kreisschreiben 5.3 «Besondere Versicherungsformen mit eingeschränkter Wahl der Leistungserbringer: Nachweis Kostenunterschiede», Inkrafttreten 1. Juni 2019, vom 5. April 2019.
- Fachliteratur: siehe Referenzen im Appendix.

2.6. Algorithmus und Codierung

Die Lösungserarbeitung, wie nachfolgend beschrieben, umfasst einen dazugehörigen R-Code bestehend aus folgenden Modulen mit dem Namen:

- Auswertung_efmc_UTF8_Final_3.R
- user_input_Final.R
- 1_mc_daten_lesen_UTF8_Final_2.R

-
- 2_rueckmeldung_efmc_UTF8_Final_3.R
 - 4_run_mc_check_UTF8_Final_3.R
 - 5_plot_and_table_Final.R

Die Codierung wurde in R unter Anwendung der R-Version 3.5.3 vorgenommen. Ein lauffähiger R-Code wurde dem BAG übergeben. Der R-Code ist Teil dieses Berichts und findet sich im Anhang dieses Berichts (siehe Appendix B).

3. Diskussion des bestehenden Ansatzes

3.1. Beschreibung des bestehenden Ansatzes (KS 5.3)

Die detaillierte Beschreibung zu den Berechnungen im aktuell verwendeten Ansatz findet sich in der Beilage zum KS 5.3 (siehe Abschnitte 3 bis 5; S. 6-9). Im Rahmen der erhobenen Analysen werden die Versicherungsunternehmen dazu aufgefordert, für jedes zu genehmigende Versicherungsmodell ihre Versicherungsbestände im sogenannten «EF-MC» Format abzulegen. Dieses Format beinhaltet Daten zu den letzten 5 Jahren für alle vordefinierten Klassen. Diese Klassen setzen sich aus den Faktoren «Prämienregion», «Altersgruppe», «Geschlecht», «Franchise», «Spitalaufenthalt im Vorjahr» und «Tod im Analysejahr» zusammen (siehe dazu auch Tabelle 8 als Verweis auf die zur Verfügung gestellten Testdaten). Für alle definierten Klassen k sind vom Versicherungsunternehmen die folgenden Daten der Modelle, jeweils aufgeteilt nach Alternativ- und Basismodell, zu erheben:

Variable	Beschreibung
Daten des Alternativmodells (Modell mit eingeschränkter Wahl der Leistungserbringer)	
NMC_k	Anzahl Versicherte als Summe der Versicherungsmonate in Klasse k dividiert durch 12.
LMC_k	Summe der Nettoleistungen aller Versicherten in Klasse k , d.h., $LMC_k = \sum_i LMC_{k,i}$, wobei $LMC_{k,i}$ dem Total der für den Versicherten i bezahlten Nettoleistungen für Behandlungen im betrachteten Jahr entspricht (Stand bei Einreichung der Daten).
QMC_k	Summe der Quadrate der Nettoleistungen aller Versicherten in Klasse k , d.h., $QMC_k = \sum_i LMC_{k,i}^2$.
PMC_k	Summe der Prämien (ohne Prämienverluste) aller Versicherten der Klasse k . Die Prämien berücksichtigen sämtliche gewährten Rabatte, z.B., für Wahlfranchisen oder das Ruhen der Unfalldeckung.
$PMCO_k$	Summe der Prämien (ohne Prämienverluste) aller Versicherten in Klasse k analog der Prämie PMC_k aber ohne die Rabatte für die Modelle mit eingeschränkter Wahl der Leistungserbringer. Die Differenz zwischen dieser fiktiven Prämie und der Prämie PMC_k ergibt den effektiv gewährten Prämienrabatt.
Daten des Basisversicherungsmodells	
$NBase_k$	Anzahl Versicherte als Summe der Versicherungsmonate in Klasse k dividiert durch 12.
$LBase_k$	Summe der Nettoleistungen aller Versicherten in Klasse k , d.h., $LBase_k = \sum_i LBase_{k,i}$, wobei $LBase_{k,i}$ dem Total der für den Versicherten i bezahlten Nettoleistungen für Behandlungen im betrachteten Jahr entspricht (Stand bei Einreichung der Daten).
$QBase_k$	Summe der Quadrate der Nettoleistungen aller Versicherten in Klasse k , d.h., $QBase_k = \sum_i LBase_{k,i}^2$.

Tabelle 3 Definition der Databasis für die Berechnung im KS 5.3

Für die Berechnung der mittleren Kosten und der Varianzen dieser Mittelwerte werden nur die Klassen k berücksichtigt, welche **mindestens je 2 Versicherte bzw. 24 Versichertenmonate** im betrachteten Modell mit eingeschränkter Wahl der Leistungserbringer und in der Basisversicherung aufweisen. Der Parameter K bezeichnet die Anzahl der von der Auswertung berücksichtigten Klassen.

Versicherte in Modellen mit eingeschränkter Wahl der Leistungserbringer

Die Anzahl der von der Auswertung berücksichtigten Versicherten NMC , der Mittelwert ihrer Kosten A , die Varianz von A , und die Durchschnittsprämien PA und $PA0$ betragen:

$$NMC = \sum_{k=1}^K NMC_k$$

$$A = \frac{1}{NMC} \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i} = \frac{1}{NMC} \sum_{k=1}^K LMC_k \quad (1)$$

$$Var(A) = Var\left(\frac{1}{NMC} \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}\right) = \frac{1}{NMC^2} \sum_{k=1}^K NMC_k \cdot Var(LMC_{k,i}) = \frac{1}{NMC^2} \sum_{k=1}^K NMC_k \cdot \frac{QMC_k - \frac{LMC_k^2}{NMC_k}}{NMC_k - 1}$$

$$PA = \frac{1}{NMC} \sum_{k=1}^K PMC_k$$

$$PA0 = \frac{1}{NMC} \sum_{k=1}^K PMC0_k$$

Versicherte in Modellen mit eingeschränkter Wahl der Leistungserbringer, falls sie eine Basisversicherung abgeschlossen hätten

Im Folgenden werden die Kosten der Versicherten in Modellen mit eingeschränkter Wahl der Leistungserbringer abgeschätzt, falls sie eine Basisversicherung abgeschlossen hätten. Dazu werden für jede Klasse die Kosten der Basisversicherung auf die Anzahl der Versicherten in den untersuchten Modellen hochgerechnet mit der Annahme, dass die Mittelwerte und Varianzen aus der Basisversicherung übernommen werden dürfen. Damit ergeben sich die Durchschnittskosten B und die Varianz von B :

$$B = \frac{1}{NMC} \sum_{k=1}^K NMC_k \cdot \frac{LBase_k}{NBase_k} \quad (2)$$

$$Var(B) = \frac{1}{NMC^2} \sum_{k=1}^K NMC_k \cdot \frac{QBase_k - \frac{LBase_k^2}{NBase_k}}{NBase_k - 1}$$

Der Maximalrabatt folgt aus der Kosteneinsparung der letzten 5 Jahre $B - A$ und den zufälligen Streuungen (zwei Standardabweichungen). Zusätzlich wird die Teuerung $\frac{PA0(FJ)}{PA0}$ zwischen dem von den Kostennachweisen abgedeckten Zeitraum und dem Folgejahr berücksichtigt⁷. Dies setzt voraus, dass die Einsparungen und damit die zulässigen Rabatte proportional zu der Anzahl der in den Modellen versicherten Personen und den durchschnittlichen Kosten bzw. Prämien der Basisversicherten sind. Der Maximalrabatt R_{max} beträgt dann

$$R_{max} = (B - A + 2 \cdot \sqrt{Var(A) + Var(B)}) \cdot \frac{PA0(FJ)}{PA0}.$$

Prämientarife sind dann zulässig, falls die Bedingung $R(FJ) \leq R_{max}$, für $R(FJ)$ den zu genehmigenden Durchschnittsrabatt, eingehalten wird.

Bemerkung 2:

- (a) Zu beachten ist, dass es sich bei den einzelnen Variablen NMC_k und $NBase_k$ nicht um ganze Zahlen handelt, was die Beschreibung durch «Anzahl Versicherte» nahelegen würde. NMC_k und $NBase_k$ sind aufgrund ihrer Definition rationale Zahlen. Dies führt dazu, dass die Schätzer A und B zwar Durchschnittsschäden/-leistungen darstellen, aber nicht arithmetischen Mittelwerten entsprechen, da nicht durch die entsprechende Anzahl der Beobachtungen/Versicherten dividiert wird. Hinzu kommt, dass die rationalen Werte NMC_k und $NBase_k$ in den Formeln des KS 5.3 als Obergrenzen von Summationen der Form $\sum_{i=1}^{NMC_k}(\dots)$ verwendet werden (siehe z.B. Formel (1)), was rein formal nicht wohldefiniert ist. Die entsprechenden Definitionen der Variablen LMC_k , QMC_k ,

⁷ Für eine detaillierte Beschreibung und Berechnung der Terme $PA0(FJ)$ und $PA0$ verweisen wir auf das KS 5.3, S. 8.

$LBase_k$ und $QBase_k$ in Tabelle 3 bilden das Verständnis grundsätzlich richtig ab (als Summe über alle Versicherte, welche unabhängig von den Deckungsmonaten ist). Bei den Formeln auf S. 7 des KS 5.3 wird dies dann allerdings sachlich nicht entsprechend übernommen.

Die oben dargestellte Ungenauigkeit im KS 5.3 wird im Folgenden anhand eines Beispiels veranschaulicht. Angenommen in einer Klasse k sind neun Versicherte vermerkt, die alle jeweils ein halbes Jahr versichert sind und alle dieselbe Leistung erzeugt haben (d.h., $LMC_{k,i} = l > 0$, für $i = 1, \dots, 9$). Das KS 5.3 definiert die Anzahl Versicherte NMC_k (Verweise auf S. 6 im KS 5.3) wie folgt

$$NMC_k = \frac{1}{12} \cdot \sum_{i=1}^9 \text{Versichertenmonate}_{k,i} = \frac{1}{12} \cdot \sum_{i=1}^9 6 = \frac{54}{12} = 4.5 \quad ,$$

mit $\text{Versichertenmonate}_{k,i}$ die Monate, die ein Versicherter i in Klasse k versichert war. Darüber hinaus wird die Summe der Nettoleistungen aller Versicherten in Klasse k (ebenfalls Verweis auf S. 6 im KS 5.3) definiert als

$$LMC_k = \sum_{i=1}^9 LMC_{k,i} \stackrel{(1)}{=} \sum_{i=1}^9 LMC_{k,i} = \sum_{i=1}^9 l = 9 \cdot l \quad .$$

Die Gleichheit (1) folgt aus der Definition, da kein Verweis auf die «Anzahl Versicherte» vorgenommen wird, sondern auf das «Total ALLER Versicherten in Klasse k » verwiesen wird. In Folge ist die Berechnung der Mittelwerte der Kosten (Verweis auf S. 7 im KS 5.3) aus mathematischer Perspektive problematisch. Als Beispiel dient die Berechnung

$$\text{Var}(A) = \text{Var} \left(\frac{1}{NMC} \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i} \right) \quad .$$

In dieser Formel wird aus dem Summanden für die Klasse k , für welchen gemäss Definition ein Wert von $9 \cdot l$ errechnet wurde, nun

$$LMC_k = \sum_{i=1}^{NMC_k} LMC_{k,i} = \sum_{i=1}^{4.5} l \stackrel{(2)}{=} 4.5 \cdot l \quad .$$

Dies ist weder mathematisch wohldefiniert (da eine rationale Zahl als Obergrenze in der Summation benützt wird) noch kann es effektiv ausgewertet werden (siehe Gleichheit (2)), wenn die total bezahlten Nettoleistungen pro Versicherten i nicht der Einfachheit halber gleichgesetzt werden, wie in diesem Beispiel. Eine entsprechende Ungenauigkeit in der Formulierung im KS 5.3 betrifft analog die Darstellung bzw. die Beschriftung für die Variable A .

Um dies formal richtigzustellen, müsste für alle Klassen k ein zusätzlicher Index J_k für $i = 1, \dots, J_k$ eingeführt werden, wobei J_k die effektive Anzahl der Beobachtungen in Klasse k beschreibt. Eine formale Beschreibung der NMC_k würde dann beispielsweise wie folgt aussehen:

$$NMC_k = \sum_{i=1}^{J_k} \frac{\text{Versichertenmonate}_{k,i}}{12} \quad ,$$

wobei $\text{Versichertenmonate}_{k,i}$ wie vorher die Deckungsmonate der i -ten Beobachtung/Versicherte in Klasse k darstellt. Die Definition von LMC_k und QMC_k würde entsprechend dann durch

$$LMC_k^* = \sum_{i=1}^{J_k} LMC_{k,i}$$

und

$$QMC_k^* = \sum_{i=1}^{J_k} LMC_{k,i}^2$$

eindeutig.

Wird nun diese mathematische Definition in die folgenden Berechnungen für $Var(A)$ und letztlich auch für $Var(B)$ analog dem KS 5.3 übernommen, lässt sich die Varianz von A wie folgt darstellen:

$$\begin{aligned} Var(A) &= Var\left(\frac{1}{NMC} \sum_{k=1}^K LMC_k^*\right) = Var\left(\frac{1}{NMC} \sum_{k=1}^K \sum_{i=1}^{J_k} LMC_{k,i}\right) = \frac{1}{NMC^2} \sum_{k=1}^K J_k \cdot Var(LMC_{k,i}) \\ &= \frac{1}{NMC^2} \sum_{k=1}^K J_k \cdot \frac{QMC_k^* - \frac{(LMC_k^*)^2}{J_k}}{J_k - 1} . \end{aligned}$$

Diese Gleichung zur Bestimmung von $Var(A)$ ist nur noch bedingt mit der ursprünglichen Gleichung aus dem KS 5.3 vergleichbar und würde zu anderen Ergebnissen für $Var(A)$ führen⁸.

- (b) Die Einschränkung der Daten auf Klassen, bei denen ein Minimum von je 2 Versicherten (bzw. 24 Versichertenmonate) vorhanden ist, erscheint uns – gegeben der Herausforderung in Teil (a) der Bemerkung – aus einer mathematischen Perspektive eher restriktiv. Um die Varianzen im KS 5.3 berechnen zu können, muss sichergestellt sein, dass NMC_k und $NBase_k$ ungleich 1 sind, sodass die Summanden wohldefiniert bleiben. Sieht man einmal von den offenen Fragen in Teil (a) der Bemerkung ab, wären grundsätzlich auch Klassen, bei denen NMC_k und $NBase_k$ grösser 1 sind, mögliche Beobachtungen, die in die Berechnung mit einbezogen werden könnten. Vor allem bei sehr kleinen Versicherern bzw. bei sehr kleinen Modellbeständen könnten eventuell hier bereits weitere Beobachtungen genutzt werden. Eine entsprechende Sensitivitätsanalyse haben wir in den Analysen im Kapitel 5 aufgenommen.⁹

Beobachtung #1 – Präzision der Beschreibungen (Notation) im KS 5.3:

Die Beschreibungen im KS 5.3 sind an einigen Stellen unpräzise. Beispielsweise wird für die «Anzahl Versicherte» (d.h., NMC_k und $NBase_k$) im gegenwärtigen KS 5.3 eine rationale Zahl zugelassen (siehe auch Tabelle 3). In den Formeln auf S. 7 des KS 5.3 wird dieser Wert als Summenobergrenze in den Berechnungen für die Variable LMC_k verwendet. Dies ist weder mathematisch wohldefiniert (Summationen benötigen ganzzahlige Grenzen), noch kann es in verschiedenen Fällen praktisch ausgewertet werden. Hinzu kommt, dass es nicht der Definition der LMC_k auf S. 6 des KS 5.3 entspricht. Würden diese Ungenauigkeiten durch eine einheitliche und mathematisch widerspruchsfreie Definition behoben werden, dann würde der angepasste Schätzer $Var(A)$ von der im KS 5.3 (S. 7 als Verweis auf die Definition) definierten Formel abweichen. Entsprechende Anpassungen wären ebenfalls für die Formeln von B und $Var(B)$ nötig.

⁸ Eine Analyse war nicht Gegenstand des Mandats.

⁹ Dazu verweisen wir auf die Ergebnisse zum Ansatz 4 «KS erweitert» im Abschnitt 5.3.

Neben den mathematisch unpräzisen Ableitungen der Varianz-Formeln, könnte die nicht wohldefinierte Darstellung im KS 5.3 (zwischen Definitionen und Formeln) auch dazu führen, dass Versicherungsunternehmen die Beschreibung «Anzahl Versicherte» missverstehen, vom naheliegenden ganzzahligen Wert ausgehen und somit möglicherweise ungeeignete Informationen im EF-MC Datensatz erfassen und an das BAG übermitteln.

Eine durchdachte Beschreibung des Terms «Anzahl Versicherte» und präzise Umsetzung im KS 5.3 scheint ein unerlässlicher nächster Schritt zu sein. Wir empfehlen daher, die Definition der «Anzahl Versicherten» (d.h., NMC_k und $NBase_k$) in einem ersten Schritt zu überprüfen und zu evaluieren, ob eine ganzzahlige/rationale «Anzahl Versicherte» sinnvoll ist. In einem zweiten Schritt sollte dann die entsprechende Beschreibung im KS 5.3 angepasst werden. Dies beinhaltet entsprechend auch eine präzise Ableitung der Varianzschätzer $Var(A)$ und $Var(B)$.

Aus den oben genannten Ungenauigkeiten ergeben sich im Folgenden immer wieder Herausforderungen bei der Definition alternativer Ansätze als auch bei den Auswertungen der entsprechenden Tests. In den folgenden Abschnitten werden wir mit den im KS 5.3 definierten/formulierten Schätzern weiterarbeiten, als wenn diese wohldefiniert wären, um unter anderem alternative Ansätze mit dem aktuellen Vorgehen im KS 5.3 vergleichen zu können. Entsprechend unterliegen die vorgestellten Alternativansätze den gleichen Limitationen wie dies im aktuellen KS 5.3 der Fall ist. Der Beibehalt dieser Limitierung war aus unserer Optik nötig, um in den folgenden quantitativen Analysen (siehe Kapitel 5) die Vergleichbarkeit der Ansätze gewährleisten zu können und zudem auf dem bestehenden R-Codes des BAG aufbauen zu können. Wir werden allerdings im Folgenden an entsprechender Stelle auf diese beobachtete Ungenauigkeit für das Verständnis des Lesers verweisen, wo nötig.

Während der Projektarbeit an der Methode Varianz traten teilweise unplausible Ergebnisse auf. Die Ursache liegt im Referenzmodell und das BAG übernahm den Lead für die Aufklärung. Beim Referenzmodell kamen Fragen zur Definition Anzahl Versicherten auf. Für ganzjährig Versicherte ist die Anzahl Versicherter wohldefiniert, für kürzere Verweildauern hingegen nicht. Die Thematik der Definition Anzahl Versicherter war nicht Teil des Projektauftrages und wird vom BAG im Nachgang aufgegriffen.

3.2. Validierung des bestehenden Ansatzes (KS 5.3)

Der Berechnungsansatz im KS 5.3 folgt der Idee, Versicherte in Modellen mit eingeschränkter Wahl der Leistungserbringer (Durchschnitt A) mit Versicherten in Modellen mit eingeschränkter Wahl der Leistungserbringer, falls diese eine Basisversicherung abgeschlossen hätten (Durchschnitt B), zu vergleichen. Dafür wird in der Komponente B eine Skalierung pro Klasse der Nettoleistungen anhand der Anzahl Versicherten in der Basisversicherung auf das entsprechende Niveau im Alternativmodell vorgenommen. Ein Abgleich der so definierten Durchschnitte, mittels der Differenz $B - A$ scheint uns als Kern der Analyse sinnvoll.

Es ist allerdings nochmals anzumerken, dass es sich bei den in KS 5.3 definierten Schätzern A und B für die Durchschnittsleistungen nicht um arithmetische Mittelwerte handelt. Die Summe über alle Leistungen wird beispielsweise in A nicht durch die Anzahl der Beobachtungen (d.h., $J = \sum_{k=1}^K J_k$, siehe dazu auch Bemerkung 2a), sondern durch die KS 5.3-Definition der «Anzahl Versicherte» (z.B., NMC) dividiert. Gegeben den Definitionen der Variablen im KS 5.3, scheint dieses Vorgehen allerdings plausibel. Grundsätzlich sollte man annehmen können, dass Versicherte, die kein ganzes Jahr versichert waren, auch nicht Leistungen im Umfang eines Jahres generieren können. Würde man also mit der Anzahl der Beobachtungen dividieren, wie dies beim arithmetischen Mittel nötig wäre, würden Leistungen von Versicherten gleichwertig gewichtet, unabhängig davon ob der Versicherte ein Jahr oder weniger versichert war. Um dies zu korrigieren, scheint die Verwendung des NMC konsistent zu den Definitionen in Tabelle 3.

Von einem mathematischen Standpunkt aus wirft die Definition der Ausdrücke A und B in der Form des KS 5.3 aber Fragen auf, da viele statistische Eigenschaften (v.a. für Varianzen) nur unter Berücksichtigung von arithmetischen Mitteln zulässig sind, aber dennoch in der Form in KS 5.3 angewendet werden¹⁰. Nach Fachliteratur werden beispielsweise Varianzen der Population bzw. Stichprobenvarianzen einer Zufallsvariable X mit Beobachtungen $X_1, \dots, X_N$ wie folgt definiert:

$$\text{Var}(X) := \frac{1}{N} \sum_{i=1}^N (X_i - \bar{X})^2 \text{ und } S_X^2 := \frac{1}{N-1} \sum_{i=1}^N (X_i - \bar{X})^2 ,$$

mit $\bar{X} = \frac{1}{N} \sum_{i=1}^N X_i$.

Die Varianz des Schätzers $\bar{X}$ lässt sich dann unter der Annahme von unabhängigen und identisch verteilten X_i wie folgt schreiben:

$$\text{Var}(\bar{X}) = \frac{1}{N} \cdot S_X^2 \quad (3) .$$

Für einen Beweis verweisen wir auf Teil 1 des Beweises von Theorem 2 in Appendix C.

Da aber nun A und B nicht als arithmetische Mittelwerte definiert wurden, entsprechen $\text{Var}(A)$ und $\text{Var}(B)$ auch nicht den klassischen Varianz-Definitionen nach statistischen Standards. Daraus folgt, dass die Eigenschaft, wie in Gleichung (3) erwähnt, nicht zwingend gültig ist. Diese werden aber dennoch für die Herleitungen im KS 5.3 zugrunde gelegt.

So werden beispielsweise, um die Unsicherheit in den Schätzungen zu berücksichtigen, die Varianzen der Schätzer A und B (aus dem KS 5.3) in die Berechnung des zulässigen Maximalrabatts $R_{\max}$ integriert. Der Summand, der zur Anwendung kommt, ist

$$2 \cdot \sqrt{\text{Var}(A) + \text{Var}(B)}$$

was einer zufälligen Streuung um zwei Standardabweichungen entsprechen soll. Hier sind vier Anmerkungen zu machen:

- (a) Um die Unschärfe des Schätzers $B - A$ zu berücksichtigen, wurden verschiedene Näherungen der Varianz vorgenommen. Zum einen wird im KS 5.3 die $\text{Var}(B - A)$ durch eine obere Schranke abgeschätzt, d.h.,

$$\text{Var}(B - A) \leq \text{Var}(A) + \text{Var}(B),$$

siehe dazu auch KS 5.3, auf S. 8. Unter der Annahme, dass es sich bei $\text{Var}(A)$ und $\text{Var}(B)$ jeweils um Varianzen gemäss statistischer Definition in der Fachliteratur handelt, ist die Abschätzung aus unserer Sicht formal korrekt und nachvollziehbar, da die Berechnung der exakten $\text{Var}(B - A)$ deutlich komplexer wäre und die entsprechenden Kovarianzen von B und A berücksichtigt werden müssten. Um dies adäquat im Prozess abdecken zu können, müssten erhebliche Zusatzinformationen (auf Versicherten-Ebene) von den Versicherungsunternehmen eingefordert werden. Dieser Zusatzaufwand scheint nicht praktikabel zu sein.

- (b) Im KS 5.3 werden die Varianzen der Schätzer A und B auf Basis von Varianz-Rechenregeln abgeleitet, bei denen für die korrekte Anwendung eigentlich bedingte Annahmen einfließen müssten. Diese sind in der aktuellen Beschreibung allerdings nicht eindeutig berücksichtigt. Zum Beispiel, wenn die Herleitungen auf S. 7 im KS 5.3 detaillierter ausgeführt werden, erhält man:

¹⁰ Bei den Variablen A und B handelt es sich um kein arithmetisches Mittel. Es kann daher nicht unmittelbar gefolgert werden, dass gewisse wohlbekannte statistische Eigenschaften erfüllt sind, welche für arithmetische Mittel gegeben sind. Jedoch kann ohne eine vertiefte Analyse auch keine Aussage getroffen werden, dass diese **nicht** erfüllt sind. So müssten die statistischen Eigenschaften des im KS 5.3 verwendeten Schätzers (z.B. Erwartungstreue) untersucht werden.

$$\begin{aligned}
Var(A) &= Var\left(\frac{1}{NMC} \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}\right) \stackrel{(1)}{=} \frac{1}{NMC^2} Var\left(\sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}\right) = \frac{1}{NMC^2} \sum_{k=1}^K Var\left(\sum_{i=1}^{NMC_k} LMC_{k,i}\right) \\
&\stackrel{(1)}{=} \frac{1}{NMC^2} \sum_{k=1}^K \sum_{i=1}^{NMC_k} Var(LMC_{k,i}) \stackrel{(2)}{=} \frac{1}{NMC^2} \sum_{k=1}^K NMC_k \cdot Var(LMC_{k,i}) .
\end{aligned}$$

Die Gleichheit (1) und (2) sind allerdings nur dann zulässig, wenn davon ausgegangen werden kann, dass die einzelnen NMC_k , und damit auch das NMC , als Skalare bzw. deterministische Grössen interpretiert werden können. Da diese Grössen aber, analog zu den Variablen LMC_k , auf Basis der erhobenen Beobachtungen aus dem EF-MC Datensatz bestimmt werden, ist dies aus unserer Sicht nicht zwingend gegeben und daher können diese Grössen ebenfalls auch eine Zufallsvariable darstellen. Die obigen Gleichungen bleiben jedoch gültig, wenn implizit die Verwendung von bedingten Varianzen angenommen werden kann, welche folgerichtig lautet: $Var(A) = Var(A | NMC_k, \text{für alle } k \text{ in } 1, \dots, K)$ und somit die Varianz bzw. Schwankungen innert der Klasse k misst. Wir gehen davon aus, dass in diesen Umformungen die bedingten Varianzen so verwendet werden sollen¹¹ und die Variabilität des Schätzers A daher als Funktional der Variabilität der $LMC_{k,i}$ geschrieben werden kann (analoge Überlegungen für $Var(B)$). Generell wird im Formelwerk des KS 5.3 implizit angenommen, dass die Variablen $LMC_{k,i}$ unabhängig und innert einer Klasse k identisch verteilt sind. Davon kann aber ausgegangen werden, wenn es sich bei den $LMC_{k,i}$ um Leistungen unterschiedlicher Versicherten i handelt, welche innert einer Klasse k eine ähnliche Streuung aufweisen. Eine entsprechende Annahme ist allerdings im KS 5.3 nicht explizit ausgewiesen. In der bedingten Form als auch mittels der Annahme, dass die $LMC_{k,i}$ unabhängig und identisch verteilt sind, ist die Berechnung von $Var(A)$ dann nachvollziehbar.

- (c) Bei $Var(B)$ ist noch zu berücksichtigen, dass analog zur Berechnung der $Var(A)$ die Endformel übernommen wurde. Eine entsprechende Gewichtung der einzelnen $LBase_k$, wie sie beispielsweise in der Formel für B vorhanden ist, geht in der Varianz verloren. Würde die Varianz analog zum Ansatz bei $Var(A)$ berechnet (inkl. der Annahmen, dass es sich wieder um bedingte Varianzen handelt), würde gelten:

$$Var(B) = \frac{1}{NMC^2} \sum_{k=1}^K NMC_k \cdot \frac{NMC_k}{NBase_k} \cdot \frac{QBase_k - \frac{LBase_k^2}{NBase_k}}{NBase_k - 1} .$$

Die Berechnung der $Var(B)$, wie im KS 5.3 dargestellt, wendet keine wohlbekannten statistischen Aussagen oder Rechengrundlagen an. Dies ändert sich in den später vorgestellten Ansätzen in denen dies behoben wird.

Der Unterschied zwischen der oben abgeleiteten $Var(B)$ und der im KS 5.3 liegt im Faktor

$$\frac{NMC_k}{NBase_k} .$$

Verhalten sich die Grössen NMC_k und $NBase_k$ für viele Klassen sehr ähnlich, dann liegt dieser Faktor meist nahe bei 1 und sollte damit nur zu geringfügigen Abweichungen in den berechneten Ergebnissen für die obige Darstellung von $Var(B)$ und der entsprechenden Definition von $Var(B)$ gemäss dem gegenwärtigen KS 5.3 führen. Datenanalysen zu den uns zur Verfügung gestellten Testdaten (siehe auch Abschnitt 5.2.3) zeigen aber deutlich, dass dies in den meisten Beispielen nicht zwingend der Fall ist, sodass durchaus Abweichungen in den Ergebnissen zwischen den beiden Formeln für $Var(B)$ zu erwarten sind. Eine detaillierte Quantifizierung ist allerdings nicht erfolgt, da die Ergebnisse stark bestandsabhängig sind.

¹¹ Dies bedeutet, dass die zufällige Verteilung der Versicherten i über die Klassen k (bzw. die Schwankungen zwischen den Beobachtungen der Klassen k) keinen Einfluss auf die Bestimmung der Höhe des Maximalrabatts R_{max} haben soll.

- (d) Nimmt man abschliessend an, dass eine Zufallsvariable X mit identisch und unabhängig verteilten Beobachtungen $X_1, \dots, X_N$ einer Normalverteilung $N(\mu_X, \sigma_X^2)$ folgt, dann liegt üblicherweise in 95% der Fälle eine Beobachtung X_i , innerhalb von zwei Mal (genauer 1.96) der Standardabweichung σ_X um den Mittelwert μ_X . Das heisst, $X_i \in \mu_X \pm 2 \cdot \sqrt{\sigma_X^2}$ bzw. $\mu_X \in X_i \pm 2 \cdot \sqrt{\sigma_X^2}$ in 95% der Fälle. Für die im KS 5.3 definierte Differenz $B - A$ und unter der Annahme, dass diese Variable $B - A$ einer Normalverteilung $N(\mu_{B-A}, \sigma_{B-A}^2)$ folgt als auch unter Verwendung der Abschätzung in Teil (a) und eines Schätzers $Var(B - A)$ für σ_{B-A}^2 , dann erhält man

$$\mu_{B-A} \in (B - A) \pm z_{0.975} \cdot \sqrt{\sigma_{B-A}^2}.$$

Somit folgt für die obere Schranke des zweiseitigen Intervalls um (das unbekannte) μ_{B-A} und einsetzen der jeweiligen Schätzer und der Abschätzung aus a) oben:

$$\begin{aligned} (B - A) + 2 \cdot \sqrt{\sigma_{B-A}^2} &\approx (B - A) + 2 \cdot \sqrt{Var(B - A)} \\ &\leq (B - A) + 2 \cdot \sqrt{Var(A) + Var(B)}, \end{aligned}$$

mit $Var(A)$ und $Var(B)$ wie oben erwähnt. Da es sich dabei jeweils bereits um Schätzer für die Varianz vom Schätzer A und B handelt, wird keine weitere Division um die Wurzel aus der Anzahl der zugrundeliegenden Beobachtungen für A bzw. B vorgenommen¹². Unter der Annahme, dass $(B - A)$ näherungsweise normalverteilt ist, ist die Verwendung des Faktors 2 plausibel, wenn mehr als 30 Beobachtungen¹³ zugrunde liegen. Die Wahl des Konfidenzniveaus in Höhe von 95% ist im Zusammenhang mit dem angestrebten Vertrauensniveau zu beurteilen. Die Höhe des Konfidenzniveaus als auch die Verwendung einer derartigen Unschärfedefinition ist gemäss unserer Einschätzung eine sehr übliche Vorgehensweise in der Versicherungswirtschaft und aus unserer Sicht nachvollziehbar. Eine entsprechende Bestätigung, dass dies wirklich die ursprüngliche Intention bei der Definition des KS 5.3 war, ist nicht vermerkt oder uns bekannt.

Bei den im Folgenden dargestellten alternativen Ansätzen, werden wir daher andere Ansätze zur Bestimmung der Varianz definieren. Auf Grund der bisherigen Überlegungen in diesem Abschnitt wird die Berechnungsformel zur Abschätzung des Maximalrabatts R_{max} , wie auf Seite 13 in diesem Bericht aus dem KS 5.3 zitiert, aber beibehalten. Das bezieht auch den Faktor 2 mit ein.

Einige der oben erwähnten Einschränkungen und die Verwendung von Näherungen sind nicht nötig, wenn die Schätzer A und B im KS 5.3 in Form arithmetischer Mittelwerte definiert werden und die entsprechenden Varianzen dieser Schätzer (d.h., $Var(A)$ und $Var(B)$) auf Basis von statistischen (Standard-)Definitionen in der Fachliteratur abgeleitet würden. Wir empfehlen deshalb:

Beobachtung #2 – Definition der Schätzer A und B im KS 5.3 als (gewichtete) arithmetische Mittelwerte

Die Definition der Durchschnittsleistungen A und B beruht im gegenwärtigen KS 5.3 nicht auf arithmetischen Mitteln. Dies hat zur Folge, dass die resultierenden Varianzen der Schätzer A und B nicht unmittelbar auf Basis von statistischen Definitionen in der Fachliteratur abgeleitet werden können.

¹² Wenn eine Variable A selbst ein Durchschnittswert ist, dann folgt bereits, dass $Var(A) = \frac{s_A^2}{N_A}$ bzw. $\sqrt{Var(A)} = \frac{s_A}{\sqrt{N_A}}$.

¹³ Bei 30 Beobachtungen ist das Quantil einer t-Verteilung mit Freiheitsgrad $\nu = 30$ zum Signifikanzniveau $\frac{\alpha}{2} = 1 - \frac{0.95}{2} = 0.025$ (zweiseitig) mit 2.042 bereits nahe bei 2 und damit nahe am Quantil einer Standardnormalverteilung zum Signifikanzniveau $\frac{\alpha}{2} = 0.025$. Siehe zum Beispiel Fachliteratur Casella und Berger, 2002.

Entsprechende Eigenschaften von Mittelwerten und Varianzen sind somit nicht immer allgemein gültig, auch wenn diese im KS 5.3 generell zur Anwendung kommen.

Um eine statistisch fundierte Herleitung zu gewährleisten, empfehlen wir die aktuelle Definition der Durchschnittsleistungen im KS 5.3 zu überdenken und ggfs. durch (gewichtete) arithmetische Mittelwerte zu ersetzen. So kann sichergestellt werden, dass die resultierenden Varianzen wohldefiniert sind und auch die statistisch bekannten Eigenschaften für die verwendeten Schätzer erhalten bleiben.

Beispiel:

Wird der Schätzer A als (gewichteter) arithmetischer Mittelwert aller Beobachtungen $LMC_{k,i}$ definiert, lässt sich dieser als auch die Varianz von A (gilt analog für B und $Var(B)$) in der folgenden Form darstellen:

$$A := \frac{1}{J} \sum_{k=1}^K LMC_k = \frac{1}{J} \sum_{k=1}^K \sum_{i=1}^{J_k} LMC_{k,i} = \frac{1}{J} \sum_{k=1}^K \frac{J_k}{J_k} \sum_{i=1}^{J_k} LMC_{k,i} = \frac{1}{J} \sum_{k=1}^K J_k \sum_{i=1}^{J_k} \frac{LMC_{k,i}}{J_k},$$

wobei die rechte Seite der Gleichung den Namen «gewichtetes arithmetisches Mittel» nahelegt. Somit folgt

$$\begin{aligned} Var(A) &= Var\left(\frac{1}{J} \sum_{k=1}^K \sum_{i=1}^{J_k} LMC_{k,i}\right) \stackrel{(*)}{=} \frac{1}{J^2} \sum_{k=1}^K Var\left(\sum_{i=1}^{J_k} LMC_{k,i}\right) \\ &\stackrel{(*)}{=} \frac{1}{J^2} \sum_{k=1}^K \sum_{i=1}^{J_k} Var(LMC_{k,i}) \stackrel{(*)}{=} \frac{1}{J^2} \sum_{k=1}^K J_k \cdot Var(LMC_{k,i}) \end{aligned}$$

mit Index $i = 1, \dots, J_k$, wobei J_k die effektive Anzahl der Beobachtungen in Klasse k beschreibt und $J = \sum_{k=1}^K J_k$. Es wird ersichtlich, dass es sich bei $Var(A)$ um eine Gewichtung der Klassen-internen Varianzen $J_k \cdot Var(LMC_{k,i})$ handelt.

Zu beachten ist allerdings, dass die Gleichheit $(*)$ nur unter der Annahme erfüllt ist, dass die Variablen $LMC_{k,i}$ unabhängig und identisch verteilt sind. Davon kann aber ausgegangen werden, wenn es sich bei den Variablen $LMC_{k,i}$ um Leistungen von unterschiedlichen Versicherten handelt. Aus der Definition von A geht aus der rechten Seite der Gleichung hervor, dass eine durchschnittliche Leistung pro Versicherten i in Klasse k verwendet wird. Soll dabei noch in Analogie zur KS 5.3-Definition der Durchschnittsleistung eine Unterjährigkeit in den Daten berücksichtigen bzw. die zeitliche Verweildauer des Versicherten i (d.h., dass ein Versicherter keine 12 Deckungsmonate aufweist und entsprechend auch keine Leistungen auf 12-Monatsbasis erzeugt hat), ist in Betracht zu ziehen, ob zusätzlich eine Skalierung der $LMC_{k,i}$ auf 12 Monate erfolgen soll.

Abschliessend erfolgt in der Berechnung des Maximalrabattes R_{max} noch eine Skalierung um den Faktor

$$\frac{PA0(FJ)}{PA0}.$$

Dabei handelt es sich um eine Anpassung des R_{max} an die Teuerung zwischen dem von den Kostennachweisen abgedeckten Zeitraum (historische Daten auf der die Berechnung beruht) und dem Folgejahr (für welches die Berechnung durchgeführt werden soll). Diese Ergänzung ist aus unserer Sicht sinnvoll und nachvollziehbar. Im Folgenden werden wir diesen Aspekt aber immer ausklammern und das R_{max} immer zur einfacheren Handhabung ohne diesen Zusatzfaktor darstellen¹⁴.

¹⁴ Dieser Faktor ist nicht Gegenstand der vorliegenden Analyse.

4. Entwicklung alternativer Ansätze

4.1. Pooled- und Totale-Varianzschätzer

Abschnitt 3.1 zeigt, dass die Berechnung im Kreisschreiben 5.3 für das R_{max} einer Differenz aus zwei Durchschnittswerten folgt, wobei die einzelnen Durchschnittswerte jeweils eine Doppelsumme enthalten, und zwar: Summe über verschiedene Klassen, wobei die jeweiligen Summanden eine Summe über Nettoleistungen (ohne und mit Gewichtung in der Basisversicherung) darstellen. Diese Definition erinnert sehr stark an Grundlagen, welche beim statistischen Test der ANOVA (= «Analysis of Variance») verwendet werden. Das zugrundeliegende Konzept bzw. technische Grundlagen einer ANOVA wollen wir bei den alternativen Berechnungsansätzen anwenden. Daher wird entsprechend der Herleitung in den folgenden Abschnitten weiter ausgeführt. Es ist allerdings zu bemerken, dass wir keine ANOVA selbst ausführen werden oder deren Annahmen voraussetzen, sondern lediglich die zugrundeliegenden Konzepte übernehmen.

Für die Herleitung eines alternativen akademischen Ansatzes definieren wir einen Schätzer $\bar{Y}$ als den Mittelwert der Beobachtungen $y_{1,1}, y_{1,2}, \dots, y_{1,n_1}, y_{2,1}, \dots \geq 0$ verteilt über verschiedene Klassen K :

$$\bar{Y} = \frac{1}{N} \sum_{i=1}^K \sum_{j=1}^{n_i} y_{i,j}$$

mit

$$N = \sum_{i=1}^K n_i.$$

Die folgende Zerlegung der quadratischen Abweichungen der $y_{i,j}$ von $\bar{Y}$ ist eine wichtige Basis für die ANOVA und soll hier auch als Grundlage dienen.

Theorem 1 – Zerlegung von Quadratsummen:

Für jede Beobachtung $y_{i,j}$, mit $i = 1, \dots, K$, und $j = 1, \dots, n_i$, gilt

$$\sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{Y})^2 = \sum_{i=1}^K n_i \cdot (\bar{y}_i - \bar{Y})^2 + \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i)^2,$$

mit

$$\bar{y}_i = \frac{1}{n_i} \sum_{j=1}^{n_i} y_{i,j},$$

$$\bar{Y} = \frac{\sum_{i=1}^K n_i \cdot \bar{y}_i}{\sum_{i=1}^K n_i} = \bar{Y}.$$

Beweis Theorem 1: siehe Appendix C.

Die Summen in der Darstellung von Theorem 1 werden in der Fachliteratur (z.B. Casella und Berger, 2002, S. 536) auch als «Sum of Squares» bezeichnet. Die Gleichung in Theorem 1 zeigt deutlich, wie sich die Gesamtabweichung der Beobachtungen vom Mittelwert («Sum of Squares Total» = SST) auf die

Abweichung der individuellen Mittelwerte einer Gruppe vom Gesamtmittelwert («Sum of Squares between Groups» = SSG) als auch die Residual-Abweichung aufteilen lässt. Letztere wird auch immer wieder als die Variation innerhalb einer Gruppe oder auch Pooled Varianz bezeichnet und misst die gemeinsame Abweichung der K Gruppen von den Mittelwerten der Klasse i . Oft wird sie auch als Abweichung des Fehlerterms oder «Sum of Squares Error» = SSE bezeichnet. Theorem 1 lässt sich demnach auch in der folgenden Form schreiben

$$SST = SSG + SSE$$

Da es sich bei den hier definierten *Sum of Squares* allerdings noch nicht um (Stichproben-)Varianzen handelt, ist noch eine Skalierung/Gewichtung der Summe der quadrierten Abweichungen notwendig. Dies wird in der Literatur wie folgt definiert und führt dann zu den sogenannten mittleren Abweichungen (siehe dazu auch Casella und Berger, 2002, S. 538):

Quelle der Abweichung	Teiler	Summe von Quadraten	Mittlere Abweichung
Zwischen den Gruppen	$K - 1$	$SSG = \sum_{i=1}^K \sum_{j=1}^{n_i} n_i (\bar{y}_i - \bar{y})^2$	$MSG = \frac{1}{K-1} SSG$
Innerhalb der Gruppen/Residual	$N - K$	$SSE = \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i)^2$	$MSE = \frac{1}{N-K} SSE$
Total	$N - 1$	$SST = \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y})^2$	$MST = \frac{1}{N-1} SST$

Tabelle 4 Definition Sum of Squares und Teiler für die Bestimmung Mittlerer Abweichungen

Die entsprechenden Teiler werden in der Tabelle 4 aufgeführt, da diese bei der Definition der Mittleren Abweichung bzw. der alternativen Schätzer entsprechend berücksichtigt werden sollten, um statistische Eigenschaften, wie etwa die «Erwartungstreue¹⁵», herbeiführen zu können.

Eine entsprechende Übertragung zur Vorgehensweise bei der Zerlegung von Summen von Quadraten oder Rechenumformungen wie in Tabelle 4 dargestellt, lässt sich auf die Definition von A und B im Kreisschreiben für die Entwicklung alternativer Varianz-Definitionen vornehmen:

Theorem 2 – Pooled Varianz:

Für die Schätzer A und B aus dem KS 5.3, den Formeln (1) und (2) oben, lässt sich somit deren **Pooled Varianz-Schätzer** wie folgt definieren:

$$Var(A) = \frac{1}{NMC \cdot (NMC - K)} \cdot \sum_{k=1}^K \left(QMC_k - \frac{LMC_k^2}{NMC_k} \right)$$

und

$$Var(B) = \frac{1}{NBase \cdot (NBase - K)} \cdot \sum_{k=1}^K c_k^2 \cdot \left(QBase_k - \frac{LBase_k^2}{NBase_k} \right)$$

mit

$$c_k = a_k \cdot b = \frac{NMC_k}{NBase_k} \cdot \frac{NBase}{NMC}$$

¹⁵ Für Definition verweisen wir auf zum Beispiel auf Casella und Berger, 2002, oder die statistische Standardliteratur.

Beweis Theorem 2: siehe Appendix C.

Bemerkung 3:

- (a) Theorem 1 gilt vor allem vor dem Hintergrund, dass es sich bei den definierten Mittelwerten (z.B. $\bar{y}$ und $\bar{y}_i$) um die Mittelwertformeln nach klassischer Definition handelt (d.h., Summe der Beobachtungen dividiert durch die Anzahl der Beobachtungen). Wie in Bemerkung 2a festgehalten, handelt es sich in den Definitionen für Schätzer A und B im KS 5.3 zwar um Durchschnittsleistungen, aber eben nicht im klassischen Sinne. Diesen Aspekt gilt es im Folgenden zu berücksichtigen, da es an verschiedenen Stellen später Anpassungen nötig macht.
- (b) Grundsätzlich kann die Pooled Varianz für alle Klassen, bei denen $NMC_k > 0$ (respective $NBase_k > 0$) gegeben ist, berechnet werden. Eine Einschränkung an die Daten, damit die formale Berechnung möglich/definiert ist, gibt es im Vergleich zu Varianz im KS 5.3 aus einer theoretischen Perspektive nicht. Allerdings werden wir die Klassen dennoch auf die folgenden Kriterien einschränken $NMC_k \geq 1$ und $NBase_k \geq 1$. Für NMC_k und $NBase_k$ kleiner als 1 (was nach Definition des KS 5.3 möglich ist; siehe Teil (a) dieser Bemerkung als auch Bemerkung 2a) kommt es in vielen Fällen zu $\left(QMC_k - \frac{LMC_k^2}{NMC_k}\right) < 0$ oder $\left(QBase_k - \frac{LBase_k^2}{NBase_k}\right) < 0$, was zu einer Verringerung der Varianz führen würde. Ist beispielsweise nur ein Versicherter mit nur einem Deckungsmonat in einer Klasse k enthalten (d.h. $NMC_k = \frac{1}{12}$ und $QMC_k = LMC_k^2$), wird aus dem Summanden für die Klasse k

$$\left(QMC_k - \frac{LMC_k^2}{NMC_k}\right) = \left(QMC_k - \frac{LMC_k^2}{1/12}\right) = LMC_k^2 - 12 \cdot LMC_k^2 = -11 \cdot LMC_k^2$$

Der Summand wird demnach nicht nur negativ, sondern erfährt – unabhängig von der Höhe des entsprechenden LMC_k – auch noch eine signifikante Verstärkung. Diese hat in einigen praktischen Tests – vor allem für sehr kleine Versicherungsunternehmen bzw. kleine Modellbestände – in der quantitativen Bewertung dazu geführt, dass sogar die Pooled Varianz selbst negativ wurde. Für diese Fälle war die Varianz undefiniert, sodass kein R_{max} – wie in Abschnitt 3.1 beschrieben – berechnet werden konnte.

Im Folgenden beziehen wir uns aufgrund der Anschaulichkeit nur noch auf Verweise im Alternativmodell. Die erarbeiteten Erkenntnisse als auch Konzepte sind aber ohne Einschränkungen auf die Daten und Schätzer in der Basisversicherung übertragbar.

- (c) Teil (b) wurde mit PuS besprochen und es wurden verschiedene Optionen evaluiert:
- i. **Restriktion der Daten auf Klassen bei denen mindestens $NMC_k \geq 1$ Versicherte enthalten sind:** Ein Vorteil dieses Ansatzes ist, dass er sich leicht auf den bestehenden EM-MC Daten implementieren lässt. Wenn die Daten auf $LMC_k \geq 1$ eingeschränkt werden, wird das Risiko, dass einzelne Summanden $\left(QMC_k - \frac{LMC_k^2}{NMC_k}\right)$ extreme Negativwerte annehmen, deutlich reduziert. Dies soll verhindern, dass auch die Pooled Varianz als Ganzes negativ wird (siehe z.B. Teil (b) der Bemerkung). Durch die Einschränkung wird allerdings nicht verhindert, dass Einzeltermine dennoch negativ werden können. Diese Einzelsummanden werden dann aber in der Aggregation der Pooled Varianz-Berechnung mit vornehmlich positiven Werten ausgeglichen. In den quantitativen Analysen (siehe Kapitel 5) wird dann bestätigt, dass es insgesamt zu keinen negativen Pooled Varianz-Schätzern mehr kam. Nachteil des Ansatzes ist allerdings, dass weiterhin verfügbare Daten nicht mit in die Berechnung einbezogen werden. Dieser Ansatz wurde in Kapitel 5 implementiert.

- ii. **Erhebung zusätzlicher Daten im EF-MC Datensatz:** Auf rein theoretischer Basis (Theorem 1) sollten negative Varianzen nicht als Resultat möglich sein. Dass dies in der Praxis hier passiert, ist ein Resultat der Definition im KS 5.3 bzgl. «Anzahl Versicherte» (siehe Bemerkung 2a). Im Beweis zum Theorem 2 (siehe Appendix C) lässt sich $Var(A)$ – vor Umformungen – wie folgt darstellen:

$$Var(A) = \frac{1}{NMC} \cdot \frac{1}{NMC - K} \cdot SSE_A$$

mit

$$SSE_A = \sum_{k=1}^K \sum_{i=1}^{NMC_k} \left(LMC_{k,i} - \frac{LMC_k}{NMC_k} \right)^2$$

Aufgrund der quadrierten Summanden, kann auch hier die Varianz nicht negativ werden. Diese Formel lässt sich allerdings nicht im EF-MC Datenformat abbilden, da hier individuelle $LMC_{k,i}$ (also Daten auf Versicherten-Ebene) nötig wären. Ein Kompromiss könnte sein, dass das EF-MC Datenformat erweitert wird, dass in einer weiteren Spalte pro Klasse k der Betrag $\sum_{i=1}^{NMC_k} \left(LMC_{k,i} - \frac{LMC_k}{NMC_k} \right)^2$ vom Versicherungsunternehmen berechnet und abgelegt wird. Infolge wird dann die Variable QMC_k unter Anwendung der Binomischen Formeln errechnet, in obigen Formeln entsprechend eingesetzt und nicht weiter erhoben (analoge Vorgehensweise für $QBase_k$). Nach Diskussion mit PuS wurde diese Möglichkeit im Rahmen dieser vorliegenden Analysen nicht weiter berücksichtigt, da die Auswertung der Summe von Quadraten u.U. erheblichen Zusatzaufwand für die Versicherer bedeutet hätte.

- iii. **Kappen von negativen Summanden:** Um negative Summanden in der Formel zu verhindern, wäre auch eine vorgelagerte Maximumfunktion in der Form

$$\max \left\{ 0; \left(QMC_k - \frac{LMC_k^2}{NMC_k} \right) \right\}$$

eine Möglichkeit. Ein Vorteil ist, dass die Implementierung auf den bestehenden EF-MC Daten einfach aufgesetzt werden kann und im Total negative Varianzen verhindert werden. Allerdings führt diese Anpassung dazu, dass für verschiedene Klassen Null-Summanden in der Summe addiert werden und die Anzahl der Versicherten NMC_k weiterhin berücksichtigt werden. Dies führt insgesamt zu einer geringeren Varianz. Daraus folgt, dass je mehr sehr kleine Klassen existieren, die Varianz des Durchschnitts geringer wird, ohne dass die Leistungen der Klassen berücksichtigt werden. Da uns dies nicht eingängig erscheint, haben wir uns entschieden, diese Anpassung vorerst nicht vorzunehmen.

- (e) Ein Spezialfall in der Betrachtung der Pooled Varianz ist, falls $NMC_k = 1$ oder $NBase_k = 1$. Dies kann auf verschiedene Fälle zurückgeführt werden:

- (i) In der Klasse k befindet sich genau ein Versicherter, der auch genau 12 Monate versichert war. In diesem Fall ist allerdings $LMC_k^2 = QMC_k$ und mit $NMC_k = 1$ wird dann der Summand $\left(QMC_k - \frac{LMC_k^2}{NMC_k} \right) = \left(QMC_k - \frac{LMC_k^2}{1} \right) = 0$. Eine Berücksichtigung dieser Klasse in der Berechnung der Varianz bedeutet also, dass die Varianz kleiner wird. Dies macht aus unserer Sicht Sinn, da mit der Pooled Varianz die Varianz innerhalb der Klassen gemessen

werden soll. Besteht eine Klasse nur aus einer Beobachtung, ist die Varianz Null, und geht in der Form auch in die aggregierte Berechnung mit ein.

- (ii) In Klasse k werden mehr als eine Person erfasst, die aber zusammen genau 12 Deckungsmonate aufweisen (z.B. Person 1 mit 4 Monaten, Person 2 mit 8 Monaten). Unter der Annahme $LMC_{k,i} \geq 0$, sollte gelten, dass auch $LMC_k^2 \geq QMC_k$ (binomische Formel). Damit wäre der Summand $\left(QMC_k - \frac{LMC_k^2}{1}\right) \leq 0$ und würde somit zu einer Reduktion der Varianz führen. Hinzu kommt eine verstärkte Reduktion in der Berechnung der Gesamtvarianz durch den grösseren Divisor. Dies scheint auf den ersten Fall nicht einleuchtend zu sein.

Da es sich bei diesem Fall allerdings um einen sehr seltenen Zustand handeln dürfte, bei dem mehrere Versicherte in einer Klasse zusammengefasst werden und genau 12 Deckungsmonate erreichen, haben wir beschlossen – auch in Hinblick auf die Komplexität des Modells – diesen Fall analog wie (i) mit zu berücksichtigen.

Theorem 3 – Totale Varianz:

Für die Schätzer A und B aus dem KS 5.3, Formeln (1) und (2), lässt sich deren **Totale Varianz-Schätzer** wie folgt definieren:

$$Var(A) = \frac{1}{NMC \cdot (NMC - 1)} \cdot \left(\sum_{k=1}^K QMC_k - \frac{(\sum_{k=1}^K LMC_k)^2}{NMC} \right)$$

und

$$Var(B) = \frac{1}{NBase \cdot (NBase - 1)} \cdot \left(\sum_{k=1}^K c_k^2 \cdot QBase_k - \frac{(\sum_{k=1}^K c_k \cdot LBase_k)^2}{NBase} \right)$$

mit

$$c_k = a_k \cdot b = \frac{NMC_k}{NBase_k} \cdot \frac{NBase}{NMC}.$$

Beweis Theorem 3: siehe Appendix C.

Bemerkung 4:

- (a) $Var(B)$ aus Theorem 3 lässt sich durch einfache Umformung auch analog schreiben als:

$$\begin{aligned} Var(B) &= \frac{1}{NBase \cdot (NBase - 1)} \cdot \left(\sum_{k=1}^K c_k^2 \cdot QBase_k - NBase \cdot B^2 \right) \\ &= \frac{1}{(NBase - 1)} \cdot \left(\frac{\sum_{k=1}^K c_k^2 \cdot QBase_k}{NBase} - B^2 \right) \end{aligned}$$

Diese Form wurde final in den R-Codes implementiert. Die Darstellung in Theorem 3 wurde gewählt um die Analogie des Ansatzes zwischen $Var(A)$ und $Var(B)$ aufzuzeigen.

- (b) Die Totale Varianz ist für alle $NMC_k > 0$ und $NBase_k > 0$ theoretisch definiert. Jedoch besteht wie beim Pooled Varianz Ansatz bei Verwendung der Definition von NMC_k und $NBase_k$ gemäss KS 5.3 die Möglichkeit, dass die Totale Varianz bzw. Einzelterme ebenfalls negativ werden können (siehe analog Bemerkung 3 b) und Bemerkung 3 e) ii) als Beispiel). Dies wäre nicht der Fall, wenn arithmetische Mittel verwendet werden, siehe Beobachtungen #1 und #2. Anpassungen bzw. das

Behandeln von Spezialfällen für den Nenner entfällt hier jedoch, da für die gesamten Kollektive davon ausgegangen werden kann, dass $NMC > 1$ bzw. $NBase > 1$. Dies bedeutet aber auch, dass mit der Totalen Varianz grundsätzlich alle zur Verfügung stehenden Daten bei der Berechnung (ohne Einschränkungen) verwendet werden können.

- (c) Die Totale Varianz lässt sich berechnen, sobald 2 Klassen (d.h., $K \geq 2$) existieren.
- (d) Die Totale Varianz berücksichtigt zudem die Variabilität zwischen den Klassen k .
- (e) Anzumerken ist, dass aufgrund der Definition und den Erkenntnissen aus Theorem 1, die Totale Varianz immer grösser sein sollte als die Pooled Varianz. Eine Verwendung der Totalen Varianz in der Berechnung des R_{max} führt auf Grund der Definition grundsätzlich zu einem grösseren Maximalrabatt als die Verwendung der Pooled Varianz als auch der Varianz-Definition gemäss KS 5.3.

4.2. Imputationsansatz

Selbst mit der Einführung der Totalen Varianz können unter Umständen nicht alle Klassen, die in den Daten enthalten sind, in den Berechnungen berücksichtigt werden. Eine Klasse k , die Versicherte im Alternativmodell (d.h., $NMC_k > 0$) zählt, in der aber in der Basisversicherung keine Versicherten (d.h., $NBase_k = 0/NA$) vorhanden sind, kann nicht in der Berechnung berücksichtigt werden (siehe dazu auch Abbildung 1 zur Veranschaulichung). Schwachstelle (2) zeigt deutlich, dass dieses Szenario in den zukünftigen Betrachtungen immer mehr an Gewicht gewinnen wird. Wir definieren hier einen Ansatz, mit welchem dieses Problem mittel- bis langfristig Abhilfe geschaffen werden könnte.

	NMC_k	LMC_k	QMC_k		$NBase_k$	$LBase_k$	$QBase_k$
k_1				k_1			
k_2				k_2			
k_3				k_3			
...				...			
...				...			

Daten verfügbar	Fehlende Daten
--------------------	-------------------

Abbildung 1 Veranschaulichung für Notwendigkeit einer Imputationslösung

Bemerkung 5:

In den uns zur Verfügung gestellten Daten kam es vor allem bei dem kleinen Versicherer (bzw. sehr kleinen Modellbeständen) oft zu dem Fall, dass Klassen k aus der Berechnung ausgeschlossen werden mussten, da zwar Daten in der Basisversicherung vorhanden waren, aber keine Alternativ-Referenz vorlag (d.h., $NMC_k = 0$ und $NBase_k > 0$). Dies erscheint sinnvoll, da ja ein Kostennachweis für das Alternativmodell und nicht die Basisversicherung (stellt hier die Referenz dar) berechnet werden soll. Dieser Punkt wurde mit PuS diskutiert und dem PuS zur Verifizierung übergeben. In den folgenden Abschnitten wird dieser Fall daher nicht weiter betrachtet oder in irgendeiner Weise durch die Imputation adressiert.

	Region	Alter	...	Tod	NBase _k	LBase _k	QBase _k
k ₁							
k ₂				A			
k ₃							
...		B					
...							

Abbildung 2 Grundlegende Idee beim Imputationsansatz

Die grundlegende Idee beim ersten Imputationsansatz wird in Abbildung 2 verdeutlicht. Ausgehend von einem EF-MC Datensatz eines Versicherers gibt es vollständig verfügbare Klassen (dargestellt in Abbildung 2 durch den Bereich A). Basierend auf diesen Daten wird mittels eines Ansatzes der Art Regression ein Zusammenhang zwischen der durchschnittlichen Leistungshöhe $LBase_k/NBase_k$ einer Klasse k mit den zugehörigen Indikatoren für die Klassendefinition (Prämienregion, Altersgruppe, Geschlecht, Franchise, Spitalaufenthalt im Vorjahr und Tod im Analysejahr; siehe dazu auch Tabelle 8) hergestellt. Ist ein Zusammenhang bestimmt (d.h., Regressionsgerade berechnet - Fitting), wird eine so-geannte Prediction mittels der unvollständigen Daten (dargestellt in Abbildung 2 durch den Bereich B) für die fehlenden Basisdaten (dargestellt in Abbildung 2 durch den roten Bereich) hergestellt.

Imputationsansatz I1:

Schritt 1 - Fitting:

Basierend auf allen Klassen k im EF-MC Datensatz bei denen $NBase_k, NMC_k > 0$ ist, wird die folgende Log-Lineare Regression durchgeführt:

$$y_k = \log\left(\frac{LBase_k}{NBase_k}\right) = \beta_0 + \beta_1 \cdot \text{Prämienregion}_k + \beta_2 \cdot \text{Altersgruppe}_k + \beta_3 \cdot \text{Geschlecht}_k \\ + \beta_4 \cdot \text{Franchise}_k + \beta_5 \cdot \text{Spitalaufenthalt im Vorjahr}_k + \beta_6 \cdot \text{Tod im Analysejahr}_k + \epsilon$$

mit den folgenden Beschreibungen:

Parameter/Variable	Beschreibung
Parameter:	
β_0	Achsenabschnitt
$\beta_1, \dots, \beta_6$	Prädiktoren für die entsprechenden unabhängigen Variablen
ϵ	Fehlerterm gemäss einer Normalverteilung; beschreibt alle Differenzen der Beobachtungen der Regressionsgeraden
Variablen:	
Prämienregion_k	Indikator-Variable/Faktor für die Prämienregionen, zu der die Daten der Klasse k gehören (siehe dazu auch Tabelle 8)
Altersgruppe_k	Indikator-Variable/Faktor für die Altersgruppe, zu der die Daten der Klasse k gehören (siehe dazu auch Tabelle 8)
Geschlecht_k	Indikator-Variable/Faktor für das Geschlecht (F/M), zu dem die Daten in Klasse k gehören (siehe dazu auch Tabelle 8)
Franchise_k	Indikator-Variable/Faktor für das Franchise-Segment (HOCH/TIEF), zu dem die Daten der Klasse k gehören (siehe dazu auch Tabelle 8)
$\text{Spitalaufenthalt im Vorjahr}_k$	Indikator-Variable/Faktor für den Flag, ob ein Spitalaufenthalt im Vorjahr (JA/NEIN) nötig war, zu dem die Daten der Klasse k gehören (siehe dazu auch Tabelle 8)

<i>Tod im Analysejahr_k</i>	Indikator-Variable/Faktor für den Flag, ob die Versicherten in den Daten der Klasse <i>k</i> im Analysejahr gestorben sind (JA/NEIN) (siehe dazu auch Tabelle 8)
---------------------------------------	--

Tabelle 5 Beschreibung Parameter/Variablen im Imputationsansatz I1

Für die Theorie zur Linearen Regression als auch die entsprechenden Berechnungsmethoden (z.B. Maximum-Likelihood) verweisen wir auf die statistische Standardliteratur (z.B., Casella und Berger, 2002). Als Resultat der Regression erhält man die Schätzer $\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_6$ für die unbekannten Parameter $\beta_0, \beta_1, \dots, \beta_6$.

Schritt 2 – Prediction:

Basierend auf allen Klassen *k* im EF-MC Datensatz bei denen $NMC_k > 0$, aber $NBase_k = 0/NA$, wird anschliessend mittels der Schätzer $\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_6$ und den entsprechenden Beobachtungen *Prämienregion_k*, ..., *Tod im Analysejahr_k* die folgende Schätzung vorgenommen:

$$\hat{y}_k = \hat{\beta}_0 + \hat{\beta}_1 \cdot \text{Prämienregion}_k + \hat{\beta}_2 \cdot \text{Altersgruppe}_k + \hat{\beta}_3 \cdot \text{Geschlecht}_k + \hat{\beta}_4 \cdot \text{Franchise}_k \\ + \hat{\beta}_5 \cdot \text{Spitalaufenthalt im Vorjahr}_k + \hat{\beta}_6 \cdot \text{Tod im Analysejahr}_k$$

Aus der Schätzung $\hat{y}_k$ lässt sich dann mit der Annahme von $NBase_k = 1$ eine Schätzung $\widehat{LBase}_k$ für die Leistung $LBase_k$ wie folgt bestimmen:

$$\widehat{LBase}_k = e^{\hat{y}_k} \cdot NBase_k.$$

Bemerkung 6:

- Die im Ansatz I1 durchgeführte Regression beinhaltet ein Logarithmieren der abhängigen Variablen $LBase_k/NBase_k$. Voraussetzung für eine Lineare Regression ist, dass der Fehlerterm ϵ einer Normalverteilung folgt. Ist die abhängige Variable allerdings unsymmetrisch (z.B. rechtsschief) und hat zusätzlich schwerere Tails als die Normalverteilung, ist die Normalität des Fehlerterms nicht zwingend sichergestellt. Diese Beobachtung kommt bei aktuariellen Betrachtungen von Schadendaten sehr oft vor. In der aktuariellen Praxis wird dann üblicherweise eine vorherige Transformation der abhängigen Variablen mit dem Logarithmus vorgenommen, um einen linearen Zusammenhang zu erhalten. Tests mittels Anwendung einer «normalen» Linearen Regression (ohne vorherige Log-Transformation) wurden geprüft und haben wie erwartet zu keinen verlässlichen Ergebnissen geführt.
- Die Transformation der abhängigen Variablen $\frac{LBase_k}{NBase_k}$ mit dem Logarithmus führt allerdings auch dazu, dass Klassen im Fitting-Datensatz (d.h., Klassen *k* bei denen $NBase_k, NMC_k > 0$) bei denen in der Basisversicherung keine Leistungen bezahlt wurden (d.h., $LBase_k = 0$), nicht in der Regression berücksichtigt werden können (Logarithmus von Null ist nicht definiert). Da diese Klassen mit vornehmlich «gesunden Versicherten» im Fitting nicht berücksichtigt werden, können grundsätzlich entsprechende Klassen auch in der Prediction nicht erzeugt werden. Dies führt dazu, dass die imputierten Leistungen der Basisversicherung tendenziell höher ausfallen dürften, was zu einem höheren Maximalrabatt R_{max} führen kann. Die Datenanalysen (siehe Abschnitt 5.3) zeigen allerdings, dass vergleichsweise wenige Klassen in der Basisversicherung mit «gesunden Versicherten» existieren. Der Einfluss auf das R_{max} ist also eher als gering einzuschätzen.
- Die Annahme in der Prediction, dass $NBase_k = 1$ ist wichtig, da so auch eine entsprechende Schätzung für $QBase_k$ gewährleistet ist. Würde eine andere Skalierung gewählt (was durchaus möglich wäre), ist die folgende Gleichung nicht mehr uneingeschränkt gegeben:

$$QBase_k = LBase_k^2.$$

Wird allerdings nur ein Versicherter (d.h., $NBase_k = 1$) in den fehlenden Daten erzeugt, müsste diese Beobachtung unter dem KS 5.3 wieder eliminiert werden, sodass eine Kombination aus Imputation und Varianzen nach KS 5.3 nicht sinnvoll ist.

Zu beachten ist allerdings auch, dass die Pooled Varianz (wenn sie auf dem Datensatz inkl. imputierten Daten berechnet wird) geringer wird, da im Fall $NBase_k = 1$ mit nur einem Versicherten der Summand der Klasse k Null wird (siehe auch Bemerkung 3d). Entsprechend sind die Ergebnisse der Imputation in Kombination mit der Pooled Varianz unter diesem Aspekt zu berücksichtigen und auch zu bewerten.

Wie Teil (b) der Bemerkung zuvor aufgezeigt hat, werden Klassen mit «gesunden Versicherten» (d.h., $NBase_k > 0$ und $LBase_k = 0$) im Fitting nicht berücksichtigt. Um diesen Punkt in einem Ansatz zu adressieren haben wir noch einen weiteren alternativen Imputationsansatz (I2) definiert.

Imputationsansatz I2:

Die Idee des Imputationsansatzes I2 beruht darauf, in einem ersten Schritt Klassen k zu erzeugen (d.h., Prediction), bei denen keine Leistungen auftraten (d.h., im Detail $\widehat{NBase}_k = 1$ und $\widehat{LBase}_k = 0$). Dies lässt sich auf eine Fragestellung «Ist ein Ereignis eingetreten oder nicht?» herunterbrechen. Für diese Art der Fragestellung eignet sich eine Logistische Regression. Die Logistische Regression beruht auf dem Prinzip, die Chance für das Eintreten eines Ereignisses zu berechnen. Ein zentraler Bestandteil dieser Diskussion ist, die sogenannte Log-Odds-Funktion («odds» Englisch für Chancen) für eine Zufallsvariable Y (nimmt den Wert 1 an, wenn das Ereignis eintritt, 0 sonst):

$$\log \left(\frac{P[Y = 1 | X = x_i]}{1 - P[Y = 1 | X = x_i]} \right) = \log \left(\frac{p}{1 - p} \right),$$

wobei $p = P[Y = 1 | X = x_i]$ die bedingte Wahrscheinlichkeit darstellt, dass das Ereignis eintritt, gegeben weiterer Informationen x_i (wobei üblicherweise x_i mittels Regressionsprädiktoren dargestellt wird – daher der Name «Logistische Regression»; siehe dazu z.B. Casella und Berger, 2002).

Sind Klassen identifiziert, in denen keine Leistungen zu erwarten sind (d.h., im Detail $\widehat{NBase}_k = 1$ und $\widehat{LBase}_k = 0$), kann für alle anderen Klassen der Imputationsansatz I1 angewendet werden. Eine direkte Modellierung von Klassen, bei denen keine Leistungen zu bezahlen sind, ist mit Ansatz I1 nicht direkt möglich, da im Fitting-Prozess $\frac{LBase_k}{NBase_k} = 0$ mit einbezogen werden müsste, was sich aufgrund der Log-Transformation in der Linearen Regression aber nicht bestimmen lässt (siehe Bemerkung 6b).

Schritt 1a – Fitting Logistische Regression:

Basierend auf allen Klassen k im EF-MC Datensatz, bei denen $NBase_k, NMC_k > 0$ ist, wird die folgende Logistische Regression durchgeführt:

$$z_k = \gamma_0 + \gamma_1 \cdot \text{Prämienregion}_k + \gamma_2 \cdot \text{Altersgruppe}_k + \gamma_3 \cdot \text{Geschlecht}_k + \gamma_4 \cdot \text{Franchise}_k \\ + \gamma_5 \cdot \text{Spitalaufenthalt im Vorjahr}_k + \gamma_6 \cdot \text{Tod im Analysejahr}_k + \epsilon$$

Die Beschreibung der Parameter und Variablen ist analog zum Ansatz I1 (siehe auch Tabelle 5). Allerdings ist die abhängige Variable z_k als die Log-Odds Funktion definiert:

$$z_k = \log \left(\frac{p_k}{1 - p_k} \right).$$

Für die Theorie zur Logistischen Regression als auch entsprechenden Berechnungsmethoden (z.B. Maximum-Likelihood) verweisen wir auch hier auf die statistische Standardliteratur (z.B., Casella und

Berger, 2002). Als Resultat der Regression erhält man die Schätzer $\hat{\gamma}_0, \hat{\gamma}_1, \dots, \hat{\gamma}_6$ für die unbekannten Parameter $\gamma_0, \gamma_1, \dots, \gamma_6$.

Schritt 1b – Prediction Logistische Regression:

Basierend auf allen Klassen k im EF-MC Datensatz bei denen $NMC_k > 0$ aber $NBase_k = 0/NA$ wird mittels der Schätzer $\hat{\gamma}_0, \hat{\gamma}_1, \dots, \hat{\gamma}_6$ und den entsprechenden Beobachtungen $Prämienregion_k, \dots, Tod\ im\ Analysejahr_k$ die folgende Schätzung vorgenommen:

$$\hat{z}_k = \hat{\gamma}_0 + \hat{\gamma}_1 \cdot Prämienregion_k + \hat{\gamma}_2 \cdot Altersgruppe_k + \hat{\gamma}_3 \cdot Geschlecht_k + \hat{\gamma}_4 \cdot Franchise_k \\ + \hat{\gamma}_5 \cdot Spitalaufenthalt\ im\ Vorjahr_k + \hat{\gamma}_6 \cdot Tod\ im\ Analysejahr_k$$

Aus der Schätzung $\hat{z}_k$ lässt sich dann durch Transformation der Log-Odds Funktion ein Schätzer $\hat{p}_k$ für p_k bestimmen unter Anwendung folgender Beziehung:

$$z_k = \log\left(\frac{p_k}{1-p_k}\right) \Leftrightarrow p_k = \frac{1}{1+e^{-z_k}}$$

Dieser Schätzer $\hat{p}_k$ stellt allerdings nur die Wahrscheinlichkeit dar, dass in Klasse k das Ereignis eintritt (z.B., dass in der Klasse k keine Leistungen zu bezahlen sind). Dabei handelt es sich noch nicht um einen Wert 0/1 der für das weitere Vorgehen benötigt wird.

Es muss also noch eine Threshold (dt.: Schwellenwert) definiert werden, sodass die Wahrscheinlichkeiten $\hat{p}_k$ über diesem Wert zu einer 1 transformiert werden können bzw. Werte darunter zu einer 0. Für diese Bestimmung müssen zusätzliche Analysen durchgeführt werden, für die ebenfalls Daten nötig sind. Deshalb muss der ursprünglich verwendeten Datensatz im Fitting in zwei Teile – den Fitting- und Test-/Validierungs-Datensatz – aufgeteilt werden. Während das Fitting aus Schritt 1a auf dem Fitting-Datensatz erfolgt, wird im Test-/Validierungs-Datensatz der Threshold bestimmt. Weitere Details zur Bestimmung des «Thresholds» finden sich z.B. in Brownlee (2020).

Schritt 2 – Modellierung der LBase_k:

Die Berechnung erfolgt hier analog zum Imputationsansatz I1, jedoch wird eine Gewichtung mit einer Indikatorfunktion als Ergebnis aus Schritt 1 vorgenommen. Es werden die $\widehat{LBase}_k$ aus I1 noch mit dem Ergebnis aus Schritt 1b multipliziert, d.h.

$$\widehat{LBase}_k^* = \widehat{LBase}_k \cdot I(\hat{p}_k \geq Threshold)$$

wobei $I(\hat{p}_k \geq Threshold)$ eine Indikatorfunktion darstellt (d.h., nimmt den Wert 1 an, wenn $\hat{p}_k \geq Threshold$ und wird 0, falls $\hat{p}_k < Threshold$).

Bemerkung 7:

- (a) Der Imputationsansatz I2 folgt einer gängigen Praxis, welche in der aktuariellen Schadenmodellierung (sogenannter Frequenz-Schadenhöhen Ansatz) Anwendung findet.
- (b) Zur besseren Darstellung, welche Daten in den einzelnen Schritten verwendet werden, verweisen wir auf die folgende Aufstellung:

Schritt		Klassen für das Fitting	Klassen für die Prediction
1	Logistische Regression	Datensatzes A (Trainingsdatensatz) in Abbildung 2 (Schritt (1a)), aufgeteilt in Fitting- und Test-/Validierungs-Datensatz	Datensatz B in Abbildung 2 (zur

			Verwendung in Schritt (1b))
2	Log-Lineare Regression (analog wie in I1)	Datensatz A aus Abbildung 2 (allerdings mit Einschränkung auf Klassen bei denen $NBase_k > 0$ und $LBase_k \neq 0$; siehe dazu auch Teil (a), aus Bemerkung 6, in dem dieser Einschränkung bereits im Detail diskutiert wurde)	Datensatz B aus Abbildung 2

Tabelle 6 Details zu verwendeten Daten im Imputationsansatz I1 und I2

- (c) Die deskriptiven Analysen auf den Daten haben gezeigt, dass vor allem für kleine Versicherer (bzw. sehr kleine Modellbestände) nur sehr wenige Klassen in der Basisversicherung mit ausschliesslich «gesunden Versicherten» (d.h., $NBase_k > 0$ und $LBase_k = 0$) existieren. Das Verhältnis von Null- zu Eins-Werten im Fitting-Datensatz ist deshalb überaus unausgewogen, was erhebliche Probleme bei der Bestimmung der Regressionsparameter verursacht und zu schlechten Ergebnissen in der Regression führt. Erschwerend kommt hinzu, dass der Trainingsdatensatz nochmals in Fitting- und Test-/Validierung-Daten unterteilt werden muss. Das Ungleichgewicht wird bei zufälliger Zuteilung der Daten zum Fitting- bzw. Test-/Validierungs-Datensatz damit noch erhöht, sodass für verschiedenste Fälle in der Berechnung (siehe Abschnitt 5.3 oder Appendix D) unrealistische Ergebnisse erzeugt werden. Die Ergebnisse unter Imputationsansatz I2 sind deshalb mit Vorsicht zu geniessen und sind aus unserer Sicht nicht für die finale Verwendung in einem revidierten KS 5.3 geeignet. Hinzu kommt, dass dieser Imputationsansatz komplex ist und auch erheblichen Aufwand in der Kommunikation erfordern dürfte.

5. Bewertung der Ansätze

5.1. Qualitative Bewertung

In den Kapiteln 3 und 4 wurden bereits verschiedenste Vor- und Nachteile des Ansatzes im KS 5.3 als auch in den alternativen Ansätzen diskutiert und vermerkt. In diesem Abschnitt werden diese nochmals aggregiert dargestellt.

Ansatz	Vorteile	Nachteile
KS 5.3	• Von der Branche grundsätzlich akzeptierter Ansatz	• Verschiedenste Ungenauigkeiten im Zusammenhang mit «Anzahl Versicherten» (Beobachtung #1) • Grosser Teil der Klassen wird in der Berechnung des Kostennachweises nicht berücksichtigt • Setzt voraus, dass die Anzahl der Versichertenmonate in einer Klasse ein Skalar ist (siehe Seite 17)
Pooled Varianz	• Fusst auf einem durchdachten mathematischen Konzept (Grundlagen zur ANOVA) • Weniger Restriktionen an die Daten als KS 5.3 • Misst eine gewichtete Varianz innerhalb der Klassen unter Verwendung von (theoretisch¹⁶) allen Klassen und ist damit der KS 5.3-Varianz am ähnlichsten, aber stabiler • Anpassungen im bestehenden Code nur gering	• Im aktuellen KS 5.3 Set-up (v.a. Definition «Anzahl Versicherte») sind auch hierfür Einschränkungen an die Daten notwendig (negative Varianz)
Totale Varianz	• Fusst auf einem durchdachten mathematischen Konzept (Grundlagen zur ANOVA) • Stabiler Schätzer für die gesamte Varianz, da diese die Variabilität im gesamten Portfolio erfasst (i.e. Klasseninterne und Zwischen-Klassen Schwankungen) • Weniger Restriktionen an die Daten als KS 5.3 und Pooled Varianz • Anpassungen im bestehenden Code nur gering • Kann berechnet werden, solange nur zwei Klassen existieren • Aus theoretischer Sicht sind auch bei diesem Ansatz Einschränkungen an die Daten notwendig (negative Varianz). In der Praxis ist die Wahrscheinlichkeit dafür jedoch gering.	• Führt in der Theorie zu einem grösseren Maximalrabatt, da die Totale Varianz grösser als Pooled Varianz (bzw. Varianz im KS 5.3) ist

¹⁶ In der Praxis sind im aktuellen KS 5.3 Set-up (siehe auch entsprechende Formulierung des Nachteils) Einschränkung an die Daten notwendig. Würden Empfehlungen #1 bzw. #2 implementiert, können auch für die Pooled Varianz alle Klassen verwendet werden und entsprechende Einschränkungen sind nicht mehr nötig.

Imputation I1	• Fusst auf einem aktuariell/statistisch anerkannten Regressionsansatz (wie auch im Risikoausgleich) • In Kombination mit Totale Varianz können alle Klassen in der Berechnung des Kostennachweises berücksichtigt werden	• Klassen, bei denen keine Leistungen aufgetreten sind (d.h., $LBase_k = 0$ und $NBase_k > 0$), können nicht modelliert werden bzw. eine Imputation vorgenommen werden. Dies führt u.U. zu höherem Maximalrabatt (Bemerkung 6) • Komplexität des Modells hoch • Eventuell schwerer zu kommunizieren und für den Einzelversicherer schwerer nachzubilden • Benötigt ggfs. grössere Anpassungen im Code (jährliche Kalibration) und ist wartungsintensiver¹⁷
Imputation I2	• Fusst auf einem aktuariell/statistisch anerkannten Modellierungsansatz • Genauere Modellierung von Klassen bei denen keine Leistungen aufgetreten sind (d.h., $LBase_k = 0$ und $NBase_k > 0$) • In Kombination mit Totale Varianz können alle Klassen in der Berechnung des Kostennachweises berücksichtigt werden	• Komplexität des Modells sehr hoch • Eventuell schwerer zu kommunizieren und für den Einzelversicherer schwerer nachzubilden • Erhebliches Ungleichgewicht in den Daten führt zu unzureichenden Modellierungsergebnissen in der Logistischen Regression • Benötigt ggfs. grössere Anpassungen im Code und ist wartungsintensiver (siehe auch analog Fussnote bei I1)

Tabelle 7 Qualitative Bewertung

Alle hier aufgeführten Ansätze sind auf dem bestehenden EF-MC Datenformat aufgesetzt. In den Diskussionen mit dem BAG wurde der Eindruck vermittelt, dass dies eine hohe Bedeutung hat. Ansätze bei denen zum Beispiel direkt auf die Versicherten-Daten (z.B. EFIND) zurückgegriffen werden müsste, sind in dem Detaillierungsgrad deshalb abschliessend nicht verfolgt worden. Nach unserem Verständnis sind entsprechende Ansätze basierend auf Versicherten-Daten für das BAG nicht präferiert, da u.U. ein erheblicher Time-Lag zwischen Datenerhebung und Auswertung bestehen könnte, der in der Praxis nicht gelöst werden kann.

Haupterkennntnis – Verwendung des alternativen Ansatzes «Totale Varianz mit Imputation I1» Gemäss unseren Analysen und erhaltenen Ergebnissen zur Messung der Leistungsfähigkeit der Schätzer zeigt sich, dass sämtliche alternative Ansätze die Schwachstelle (2) adressieren können. Das führt dazu, dass die Menge der verwendeten Datensätze in den alternativen Ansätzen im Vergleich zum Referenzmodell (teils signifikant) erhöht werden kann. Allein auf Grund der theoretischen Basis ist die Implementierung eines Ansatzes, welcher auf der Totalen Varianz basiert, sinnvoll. Dieser Ansatz erfordert die wenigsten Restriktionen an den zugrundeliegenden Daten. Wird der Ansatz nach Totaler Varianz noch mit einer Imputation kombiniert, sollte auch eine adäquate Lösung für die Schwachstelle (1) ermöglicht werden. Aus unserer Sicht wäre hier der Ansatz I1 «Log-Lineare Regression ohne vorgelagerte Logistische Regression» zu präferieren, da dieser – gegeben der Datenbasis – deutlich stabilere Resultate generiert als auch leichter gegenüber der Versicherungswirtschaft zu kommunizieren sein dürfte als der Imputationsansatz I2 («Log-Lineare Regression mit vorgelagerter Logistischer Regression»).

Im Rahmen der folgenden Diskussionen zur Schwachstelle (1) «Einführung der PCG» ist es wichtig zu verstehen, dass alle Ansätze zu teilweise erheblichen Unterschieden im Maximalrabatt führen werden.

¹⁷ Das reibungslose Funktionieren von Regressionen ist abhängig von den zugrundeliegenden Daten. Da die Daten sich jährlich ändern, sollte zumindest jährlich sichergestellt werden, dass die Regression adäquat mit den Daten arbeiten können. Hinzu kommt, dass im Idealfall die Imputation für alle Versicherer-/Modellkombinationen, für die der Maximalrabatt berechnet wird, funktionieren soll. In den weiteren Analysen (Abschnitt 5.3) wurden keine Probleme beobachtet. Aufwände für die Anpassungen an der Kalibrierung bzw. erforderliche Wartungsarbeiten sind entsprechend nicht auszuschliessen und sollten bei der Auswahl der Ansätze berücksichtigt werden.

Allgemeiner lässt sich hier festhalten, dass alle Ansätze (auch das KS 5.3) einer sehr starken Abhängigkeit von den Aggregationsmechanismen (d.h., Klassendefinition) im MF-MC Format unterliegen.

5.2. Quantitativer Ansatz zur Bewertung

Neben der qualitativen Beurteilung ist dem BAG auch eine quantitative Bewertung der Ansätze wichtig. Essentiell für diese Beurteilung ist die Ableitung objektiver Gütekriterien, anhand derer sich die verschiedenen Ansätze bewerten bzw. vergleichen lassen. Neben eher deskriptiven Beurteilungskriterien ist uns wichtig, die entsprechende Schätzunsicherheit des R_{max} unter den verschiedenen Ansätzen zu beurteilen. Um die Ungenauigkeit und Stabilität der erwähnten alternativen Schätzverfahren zu evaluieren, wird die Varianz für die Schätzung $\hat{R}_{max}$ (Statistik) bestimmt. Diese wird durch einen sogenannten Bootstrap-Ansatz ermittelt. Dabei handelt es sich um ein übliches statistisches Standardverfahren, welches für diese Art von Fragestellungen zur Anwendung kommt.

5.2.1. Bootstrap-Verfahren

Eine akademische Einführung in das Thema Bootstrapping findet sich in Efron und Tibshirani (1994). Für eine detaillierte Beschreibung, wie ein Bootstrap-/Resampling-Verfahren praktisch in R aufgesetzt werden kann, verweisen wir beispielhaft auf Wollschläger (2017).

Die Grundidee bei diesem Verfahren beruht darauf, dass man wiederholt dieselbe Statistik auf sich ändernden Daten berechnet. Dazu zieht man aus der Gesamtmenge der Daten zufällig Beobachtungen und berechnet die Statistik (in unserem Fall R_{max}) neu. Dieser Vorgang wird mehrere Male durchgeführt und anschliessend die Streuung der verschiedenen Schätzungen $\hat{R}_{max}$ berechnet. Das «zufällig Ziehen» aus der Grundgesamtheit kann man sich vorstellen wie bei einer Urne, in der die Gesamtmenge der Daten enthalten sind und aus dieser dann (entweder mit oder ohne Zurücklegen) gezogen werden kann.

Die aktuellen Berechnungen basieren alle auf einem Datensatz im Format des EF-MC, der bereits eine Aggregation von Daten auf Ebene der Versicherten (was dem EFIND-Datensatz entsprechen würde) entspricht. Deswegen, um zukünftige Konsistenz zu bewahren, wurde die Bootstrap Methode auf dieser Ebene implementiert.

Die folgende Abbildung soll den zweistufigen Prozess von EFIND zu EF-MC Aggregation verdeutlichen.

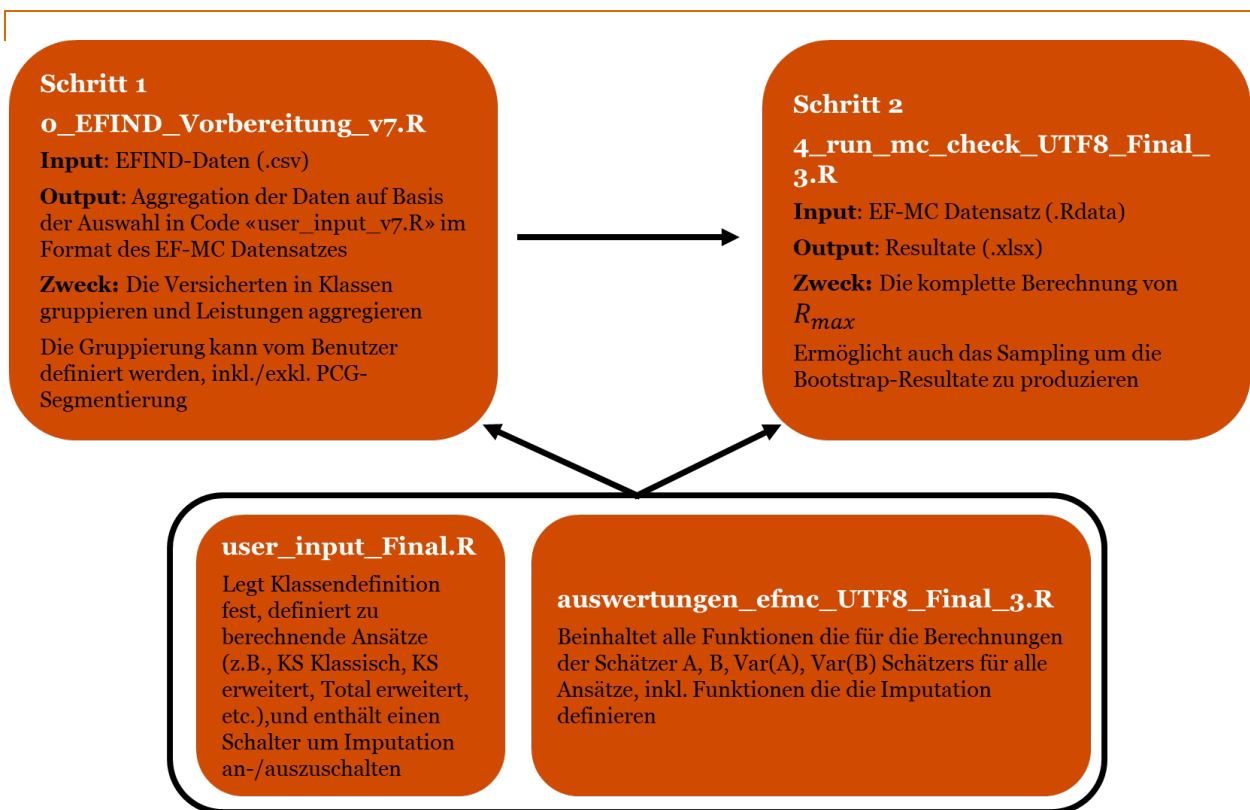


Abbildung 3 Übersicht Bootstrap-Berechnungsprozess

Die Analysen für den Bootstrap sind für die Tests in den Kapiteln 5 und 6 auf den EF-MC-Daten aufgesetzt worden. Im ersten Schritt wird die EFIND Datenbasis in das EF-MC-Format aggregiert und dient dann im zweiten Schritt als Basis für die Bootstrap Methode. Diese generiert (mittels Resampling mit Wiederholung) einen neuen Datensatz und dient zur Berechnung der zehn Analyseansätze. Dieser Prozess wird 500-Mal wiederholt, um eine stabile Verteilung von dem Schätzer $\hat{R}_{max}$ zu bekommen («4_run_mc_check_UTF8_Final_3.R»). Die Breite dieser Verteilung dient dann als Metrik, um die Analysenansätze zu vergleichen.

Für die finale Berechnung werden zwei weitere Dateien benötigt:

- «user_input_Final.R»: legt die Klassendefinition fest (z.B. ob PCG als Segmentierungskriterium einfließen sollen oder nicht), definiert welche Ansätze (siehe dazu Tabelle 10) berechnet werden sollen und enthält einen Schalter, ob eine Imputation berücksichtigt werden soll oder nicht;
- «auswertungen_efmc_UTF8_Final_3.R»: ist ein Hilfsfile, dass alle Funktionen für die Berechnung der Schätzer A, B, Var(A) und Var(B) (d.h., die technische Implementierung der Formeln und Ansätze) als auch die Umsetzung der Imputation enthält. Diese werden dann im Rahmen der Berechnungen in Schritt 2 (d.h., in File «4_run_mc_check_UTF8_Final_3.R») aufgerufen.

Als beispielhafte Darstellung dieser Methode verweisen wir auf Abbildung 4. Mit dieser Auswertung lassen sich dann Statistiken für die Unsicherheit in der Schätzung des $\hat{R}_{max}$ bestimmen.

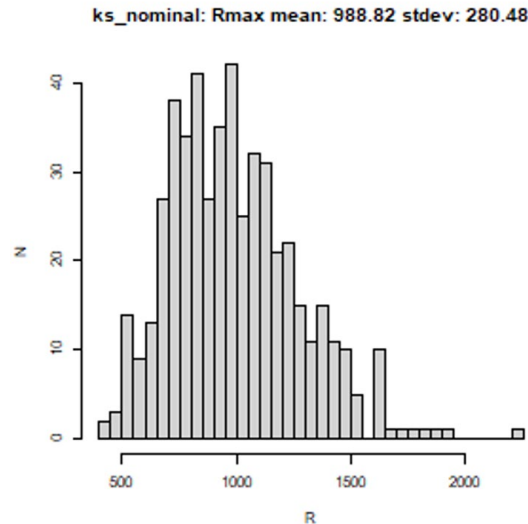


Abbildung 4 Beispielhafte Bootstrap-Visualisierung (KS klassisch)

5.2.2. Definition der Gütekriterien

Um quantitativ die Güte der alternativen Berechnungsansätze zu bestimmen, werden die folgenden Gütekriterien zur Beurteilung herangezogen und mit dem Referenzmodell verglichen:

- Anzahl der zur Berechnung verwendeten Klassen im Verhältnis aller verfügbaren Klassen (d.h., $NMC_k > 0$);
- Stabilität (d.h., Höhe/Wertigkeit) des $\hat{R}_{max}$ bei Einführung eines alternativen Ansatzes im Vergleich zum Referenzmodell. Die Frage, die hiermit verbunden ist: «Wie verändert sich die Höhe des R_{max} unter den verschiedenen Modellen bzw. kann ein stabiler Maximalrabatt beibehalten werden?»
- Varianz bzw. Standardabweichung der Schätzung $\hat{R}_{max}$ innerhalb der Bootstrap-Verteilung (d.h., Stichprobenvarianz der Bootstrap-Beobachtungen); und
- Variationskoeffizient der Schätzung $\hat{R}_{max}$ innerhalb der Bootstrap-Verteilung (d.h., als relatives Streuungsmass; Berechnung als Standardabweichung dividiert durch den Schätzer).

Im KS 5.3 werden in der Regel die Kostennachweise pro Versicherungsunternehmen als auch nach Modell separat berechnet. Dieses Vorgehen ändert sich auch im Rahmen der Alternativansätze nicht. Dies bedeutet aber auch, dass ein spezieller Modellierungsansatz für eine Krankenversicherer-Modell-Kombination objektiv gemessen vorteilhaft sein könnte, während für ein ähnliches Modell bei einem anderen Krankversicherer oder beim gleichen Krankenversicherer mit einem unterschiedlichen Modell ein anderer Modellierungsansatz vorteilhaft ist. Grundsätzliche Überlegungen über ein globales Gütekriterium (z.B. Varianz aggregiert nach Modellen) wurden mit dem PuS diskutiert, aber final nicht in der Auswertung berücksichtigt. Dies vor allem auch vor dem Hintergrund, dass nicht alle Krankenversicherer bzw. Modell-Kombinationen aus Performance-Gründen getestet werden konnten.

5.2.3. Datengrundlage

Die Datengrundlage für das KS 5.3 stellen die sogenannten EF-MC Daten dar. Für unsere Analysen haben wir einen entsprechenden Pseudo-Testdatensatz («1111-EF_MC_2021_d.xlsx») erhalten, der nach Angabe des BAG/PuS nicht repräsentativ ist. Für die Erzeugung des anonymen Testdatensatzes ist das BAG verantwortlich und die Vorgehensweise wird daher hier nicht weiter beschrieben. Dieser Datensatz enthält

die folgenden Attribute. Zusätzlich beschreiben wir unser Verständnis der entsprechenden Variablen in einer separaten Spalte:

Name der Variablen		Attributausprägungen	Beschreibung
Identifikation Kostennachweis			
Jahr	2016 – 2020	Grundsätzlich sind im MF-MC Format die Daten der letzten 5 Jahre für alle Klassen auszufüllen	
Nachweis_ID	ID1 und ID2	Für jeden eingereichten Kostennachweis ist eine eindeutige Nachweis_ID anzugeben	
Art des Modells	DIV_A und HMO_B	Beispiele für die Art des Modells (auf S. 5 des KS 5.3, sind alle möglichen Ausprägungen definiert)	
Klassen			
Prämienregion	AGo/AIo/ARo/BE1/BE2/BE3/BL1/BL2/BSo/FR1/FR2/GEo/GLo/GR1/GR2/GR3/JUo/LU1/LU2/LU3/NEo/NWo/OWo/SG1/SG2/SG3/SH1/SH2/SOo/SZo/TGo/TI1/TI2/URo/VD1/VD2/VD3/VS1/VS2/ZGo/ZH1/ZH2/ZH3	Nach KS 5.3 vorgegebene Kategorisierungsvariablen (Definition der Klassen)	
Altersgruppe	0-18/19-25/26-30/31-35/36-40/41-45/46-50/51-55/56-60/61-65/66-70/71-75/76-80/81-85/86-90/91-		
Geschlecht	F/M		
Franchise	HOCH/TIEF		
Spital Heim im Vorjahr	JA/NEIN		
Tod_im_Analysejahr	JA/NEIN		
Daten Modelle (mit eingeschränkter Wahl der Leistungserbringer)			
Anzahl Versicherte NMC	Double	Anzahl Versicherte als Summe der Versicherungsmonate in Klasse k dividiert durch 12	
Netto-Leistungen LMC	Double	Summe der Nettoleistungen aller Versicherten in Klasse k	
Leistungen im Quadrat QMC	Double	Summe der Quadrate der Nettoleistungen aller Versicherten in Klasse k	
Prämien PMC	Double	Summe der Prämien (ohne Prämienverluste) aller Versicherten der Klasse k . Die Prämien berücksichtigen sämtliche gewährten Rabatte, z.B. für Wahlfranchise oder das Ruhen der Unfallddeckung	
Prämien PMCo	Double	Summe der Prämien (ohne Prämienverluste) aller Versicherten in Klasse k analog der Prämie PMC, aber ohne die Rabatte für die Modelle mit eingeschränkter Wahl der Leistungserbringer. Die Differenz zwischen dieser fiktiven Prämie und der Prämie PMC ergibt den effektiv gewährten Prämienrabatt	
Daten Basisversicherung			
Anzahl Versicherte NBase	Double	Anzahl Versicherte als Summe der Versicherungsmonate in Klasse k dividiert durch 12	
Netto-Leistungen LBase	Double	Summe der Nettoleistungen aller Versicherten in Klasse k	
Leistungen im Quadrat QBASE	Double	Summe der Quadrate der Nettoleistungen aller Versicherten in Klasse k	

Tabelle 8 Beschreibung Attribute EF-MC Testdatensatz

Darüber hinaus haben wir einen Pseudo-Testdatensatz auf Ebene der individuellen Versicherungsnehmer erhalten, den sogenannten «EFIND»-Datensatz (Files: «BAG_Zihler_Modellrabatte_Testdata_2020_20221018_downloaded_Versand_.csv» und «BAG_Zihler_Modellrabatte_Testdata_2021_20221018_downloaded_Versand_.csv»). Auch hier sind einige uns nicht bekannte Datentransformationen durch das BAG/PuS vorgenommen worden, um die Anonymität der Daten zu gewährleisten. Die folgende Tabelle zeigt wieder die Attribute des Datensatzes inklusive unseres Verständnisses der entsprechenden Variablen:

Name der Variablen		Attributausprägungen	Beschreibung
1	BAGNR	4-stellige Zahl	Eindeutige Identifizierung des Krankenversicherungsunternehmens
2	JAHR	4-stellige Zahl	Identifizierung aus welchem Jahr die Beobachtung stammt
3	PERSONID	Zufallscode (z.B. «237486d128bc639faf307baf5031e55f03d3253f037d25264245a2f5d984a54e»)	Eindeutige Identifizierung für die versicherte Person
4	VERSTORBEN	Ja/Nein	Indikator, ob Versicherter verstorben ist oder nicht
5	LANDKANTOR	AG, ..., ZH (26 Kantons-Kennzeichnungen; bestehend aus zwei Buchstaben) + Auslands-Kennzeichnungen (z.B. DEU, FRA; bestehend aus drei Buchstaben)	Wohnort der Versicherten Person; die EU-Staaten sind für die vorliegende Analyse allerdings nicht relevant
6	REGION	Faktoren 0, 1, 2, und 3	Definition der Prämienregion innerhalb des Kantons LANDKANTON
7	GESCHLECHT	Frauen/Männer	Geschlecht der Versicherten Person
8	RISKLA	RO, R1, ...	Indikator für die Risikoklasse im Risikoausgleich, wird allerdings im Folgenden auch nicht verwendet
9	SPITALVJ	Ja/Nein/Unbekannt	Indikator ob versicherte Person im Vorjahr einen Spitalaufenthalt hatte
10-44	PCG_ABH bis PCG_DM2_hyp	0/1 pro Indikator	Insgesamt 35 Indikatoren (0/1) für verschiedenen PCG
45	TARIFTYP	BASE/DIV/HAM_RDS/HMO	Interne Spezifikation des Tariftyps (für unsere Auswertungen allerdings Verweis auf TARIFAKRO)
46	MODELLART	AND/BASE/H_DIV_A/HAM_RDS_A/ ...	Interne Spezifikation der Modellart (für unsere Auswertungen allerdings Verweis auf TARIFAKRO)
47	TARIFAKRO	Zufallscode (z.B. «5BACA8DEEC»)	Pro BAGNR, eindeutige Identifizierung für den Tarif ¹⁸
48	FRANCHISENSTUFE	Integer 1-7	Identifizierung für die Höhe der Franchise (siehe auch Spalte 49)
49	FRANCHISE	0/100/200/300/400/500/600/1000/1500/2000/2500	Höhe der Franchise
50	UNFALLEINSCHLUSS	Ja/Nein	Indikator, ob Unfalldeckung inkludiert ist oder nicht
51	BRUTTO_PRAEMIE_VERSICHERER	Double	Keine weiteren Angaben erhalten, aber für Analysen nur BRUTTO_PRAEMIE_PG relevant
52	BRUTTO_PRAEMIE_PG	Double	Prämie gemäss Prämiengenehmigung
53	NETTO_PRAEMIE	Double	Prämie abzüglich Kantonsbeitrag
54	PMCo	Double	Analog zum Wert aus dem KS 5.3, siehe dazu auch Abschnitt 3.1
55	Altersklasse_Risikoausgleich	0-18 Jahre/19-25 Jahre/26-30 Jahre/31-35 Jahre/36-40 Jahre/41-45 Jahre/46-50 Jahre/51-55 Jahre/56-60 Jahre/61-65 Jahre/66-70 Jahre/71-75 Jahre/76-80 Jahre/81-85 Jahre/86-90 Jahre/ >90 Jahre	Faktor zur Identifikation zu welcher Altersklasse (insgesamt 16) der Versicherte gehört. Die Segmentierung der Altersklassen richtet sich hier nach der im Risikoausgleich verwendeten Gruppierung. Diese sollte sich mit der aus dem EF-MC Datensatz decken
56	BRUTTOKOSTEN	Double	Effektiv angefallene Kosten eines Versicherten im Analysejahr
57	KOBE	Double	KOBE = Kostenbeteiligung; Betrag der durch Franchise und

¹⁸ Es gibt Versicherer, z.B. in einer Gruppe, welche je das gleiche Modell (gleicher TARIFAKRO) vertreiben. In diesen Fällen ist das Modell nur in Verbindung mit BAGNR eindeutig.

			Selbstbehalt durch den Versicherten zu tragen war.
58	DECKUNGSMONATE	Double	Anzahl der Monate die der Versicherte gedeckt/versicherte war (d.h., theoretisches Minimum = 0, theoretisches Maximum = 12).
59-65	Verschiedenste Datenextraktions- und Kontroll-Variablen (hier nicht näher spezifiziert da nicht relevant für die Berechnungen)		

Tabelle 9 Beschreibung Attribute EFIND Testdatensatz

Einschränkungen:

Da für uns nicht ersichtlich ist, nach welchen Massnahmen bzw. welche Transformationen in den Daten im Anonymisierungsprozess vorgenommen wurden, können wir keine Aussagen über die Repräsentativität der Daten treffen.

Allerdings sind uns folgende Widersprüche in den Daten aufgefallen:

- EF-MC: Die Schwachstelle (2) (d.h., Berechnung KS ist eingeschränkt, da Vergleichskollektive in Basisversicherung kleiner werden) ist im Testdatensatz eigentlich nicht gegeben und stellt kein Problem dar, sodass ein effektives Testen in der Vorbereitung nur bedingt möglich war;
- EFIND: verschiedene Versicherte kommen mehrfach im Datensatz vor, was sich beispielsweise durch einen Umzug in eine andere Prämienregion, einen unterjährigen Wechsel des Versicherers/der Versicherungsdeckung, oder durch die Hinzunahmen der Unfallddeckung, etc. erklären lässt. Aufgeführt sind im Testdatensatz allerdings dann Versicherte, bei denen sich die Deckungsmonate zu mehr als 12 Monaten addieren lassen. Dies betrifft etwa 20.5% der im 2020 Datensatz enthaltenen Versicherungsnehmer, bzw. 18.8% der Versicherungsnehmer im Datensatz 2021.

Hinzu kommt, dass die Nettoleistungen in diesen Einträgen gleichgesetzt sind und keine eindeutige Aufteilung der Kosten auf die Perioden möglich ist, in denen sie wirklich angefallen sind. Dies trifft auf alle Einträge zu, bei denen ein Versicherungsnehmer mehrfach im Datensatz vorkommt (ca. 29.4% der Einträge im Datensatz 2020; ca. 27.4% der Einträge im Datensatz 2021).

Basierend auf dieser Limitierung sind die uns zur Verfügung gestellten Testdaten nur zur Veranschaulichung sinnvoll gewesen, sodass die finalen Ergebnisse durch das PuS auf Echtdaten berechnet, die Ergebnisse anonymisiert und uns dann für die Auswertung zur Verfügung gestellt wurden.

Es wurden in vollem Umfang drei Krankenversicherer (KK-AJ, KK-AE und KK-AD) getestet. Nach Aussage des BAG handelt es sich dabei jeweils um einen grossen, einen kleinen und einen sehr kleinen Krankenversicherer im Schweizer Markt. Sowohl für den grossen (KK-AJ) als auch den sehr kleinen (KK-AD) Krankenversicherer wurden insgesamt zwei Modelle evaluiert. Bei dem kleinen Krankenversicherer (KK-AE) wurde nur ein Modell berechnet, sodass insgesamt fünf komplette Testauswertungen für die Analysen zur Verfügung standen. Diese werden im Folgenden als auch im Appendix D ausgewertet.

5.3. Quantitative Bewertung

5.3.1. Definition der bewerteten Ansätze/Ansatzkombinationen

In Kapitel 4 werden – im Vergleich zur Definition im KS 5.3 – zwei alternative Ansätze für die Varianzberechnung hergeleitet, die im Rahmen der hier folgenden Analysen getestet werden sollen. Beide Ansätze vermögen aber nicht den Fall zu lösen, falls es Klassen mit Modellversicherten gibt (d.h., $NMC_k > 0$), aber in der entsprechenden Klasse der Basisversicherten keine Referenz vorhanden ist (d.h., $NBase_k = 0/NA$). Der Ansatz mittels Imputation (Abschnitt 4.2) ermöglicht hier Abhilfe. Entsprechend ergeben sich auch Kombinationsansätze aus Unterschieden in der Varianzberechnung (z.B. Varianz nach KS 5.3, Pooled oder Totaler Varianz) als auch mit oder ohne Imputation (Log-Lineare Regression mit (I2) oder ohne (I1))

vorgelagerte Logistischer Regression). In allen Kombinationen erfolgt die Berechnung des R_{max} gemäss KS 5.3 mit der auf den S. 7-9 definierten Formeln. Was sich jedoch dabei ändert, ist jeweils die Definition der verwendeten Varianzschätzer $Var(A)$ und $Var(B)$ als auch die einbezogene Datenbasis. Die folgenden zehn Ansätze werden in den nächsten Abschnitten analysiert und qualitativ als auch quantitativ beurteilt:

Ansatz		Datenrestriktion		Verwendete Varianz	Verwendung Imputation	Beschreibung
		NMC_k	$NBase_k$			
(1)	KS Klassisch	≥ 2	≥ 2	Nach KS	Nein	Berechnung erfolgt gemäss KS 5.3 mit der auf S. 7-9 definierten Formeln
(2)	Pooled Klassisch	≥ 2	≥ 2	Pooled	Nein	Berechnung erfolgt gemäss KS 5.3, allerdings wird $Var(A)$ und $Var(B)$ auf S. 7 durch den Pooled Varianzschätzer (siehe Theorem 2) ersetzt. Berechnung erfolgt zur Vergleichbarkeit auf den gleichen Basisdaten wie (1)
(3)	Total Klassisch	≥ 2	≥ 2	Total	Nein	Berechnung erfolgt gemäss KS 5.3, allerdings wird $Var(A)$ und $Var(B)$ auf S. 7 durch den Totalen Varianzschätzer (siehe Theorem 3) ersetzt. Berechnung erfolgt zur Vergleichbarkeit auf den gleichen Basisdaten wie (1)
(4)	KS erweitert	> 1	> 1	Nach KS	Nein	Unter Ansatz (1) werden nur Klassen in der Berechnung berücksichtigt, bei denen sowohl im Basismodell als auch im Alternativmodell im Minimum je 2 Versicherte bzw. 24 Versichertenmonate berücksichtigt werden. Dies scheint uns sehr restriktiv und ist aus mathematischer Sicht nicht unbedingt nötig, sodass auch ein Kriterium von mehr als je 1 Versicherte bzw. 12 Versichertenmonate möglich wäre
(5)	Pooled erweitert	≥ 1	≥ 1	Pooled	Nein	Die Pooled Varianz sollte möglichst nicht für Klassen mit weniger als einem Versicherten berechnet werden (siehe dazu Bemerkung 3), sodass dieser Ansatz das Minimum für Pooled Varianz testet
(6)	Total erweitert	> 0	> 0	Total	Nein	Die Totale Varianz hat keine Restriktion an die Grösse der Klassen, sodass dieser Vorteil hier getestet wird
(7)	Pooled erweitert I1	≥ 1	≥ 1	Pooled	Ja	Hier werden in einem ersten Schritt die fehlenden Daten durch Imputation mittels Log-Linearer Regression (I1, siehe Abschnitt 0) vervollständigt und dann die Pooled (Ansatz (7) aber analog (5)) und Totale (Ansatz (8) aber analog (6)) Varianz berechnet
(8)	Total erweitert I1	> 0	> 0	Total	Ja	
(9)	Pooled erweitert I2	≥ 1	≥ 1	Pooled	Ja	Hier werden in einem ersten Schritt die fehlenden Daten durch Imputation mittels Log-Linearer Regression + vorgelagerter Logistischer Regression (I2, siehe Abschnitt 0) vervollständigt und dann die Pooled (Ansatz (9) aber analog (5)) und Totale (Ansatz (10) aber analog (6)) Varianz berechnet
(10)	Total erweitert I2	> 0	> 0	Total	Ja	

Tabelle 10 Übersicht der Ansätze im Test

5.3.1. Auswertungen

Für die zehn Ansätze werden mittels Bootstrap-Verfahren und Deskriptiver Analyse die Gütekriterien gemäss Abschnitt 5.2.2 ausgewertet. Die folgende Tabelle zeigt exemplarisch die Gütekriterien für ein ausgewähltes Versicherungsunternehmen KK-AJ mit einem Kostennachweis eines entsprechenden Modells M29.

Ansatz	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
(1) KS Klassisch	70	104	67.3%	988.82		78'669.64	280.48	28.4%
(2) Pooled Klassisch	70	104	67.3%	945.95	-4.3%	77'850.04	279.02	29.5%
(3) Total Klassisch	70	104	67.3%	995.61	0.7%	77'180.81	277.81	27.9%
(4) KS erweitert	70	104	67.3%	988.82	0.0%	78'669.64	280.48	28.4%
(5) Pooled erweitert	83	104	79.8%	958.01	-3.1%	81'500.66	285.48	29.8%
(6) Total erweitert	85	104	81.7%	1'021.48	3.3%	80'857.56	284.35	27.8%
(7) Pooled erweitert I1	97	104	93.3%	980.32	-0.9%	84'120.96	290.04	29.6%
(8) Total erweitert I1	104	104	100.0%	1'024.73	3.6%	84'982.95	291.52	28.4%
(9) Pooled erweitert I2	97	104	93.3%	934.58	-5.5%	81'442.73	285.38	30.5%
(10) Total erweitert I2	104	104	100.0%	943.18	-4.6%	77'717.30	278.78	29.6%

Tabelle 11 Quantitative Bewertung KK-AJ für Modell M29

Insgesamt wurden noch vier weitere Modelle auf Echtdaten durch das BAG getestet: Modell M30 von KK-AJ, Modell M146 von KK-AE und Modelle M117 und M150 von KK-AD. Diese zeigen ein sehr ähnliches Ergebnismuster wie die beispielhaft dargestellten Ergebnisse in Tabelle 11, sodass wir diese hier im Hauptteil nicht wiederholen. Entsprechende Auswertungen finden sich aber in Appendix D.

Die Ergebnisse in Tabelle 11 zeigen, dass Ansatz 1 – KS Klassisch eigentlich nur etwa 67% der verfügbaren Klassen im Alternativmodell in die Berechnung des Kostennachweises einfließen lässt. Das bedeutet, dass in etwa 33% der Klassen in denen Versicherte im Alternativmodell vorhanden sind, keine Versicherten bzw. weniger als 2 Versicherte bzw. 24 Deckungsmonate vorhanden sind. Ein derartig berechnetes R_{max} bezieht also etwa ein Drittel der nötigen Daten in die Berechnungen nicht mit ein und zeigt klar Schwachstelle (2) im aktuellen Ansatz des KS 5.3 auf.

Wird im KS 5.3 alleinig die Definition der Varianz ersetzt (Datenrestriktion aus (1) bleibt bestehen), zeigt sich, dass sich die Höhe des R_{max} nur geringfügig ändert (siehe dazu Ergebnisse zu den Ansätzen (2) und (3)). Der eigentliche Vorteil der Pooled und Totalen Varianz zeigt sich allerdings erst bei der Analyse in den Ansätzen (5) und (6). Es zeigt sich, dass mit der Änderung der Varianzschätzung von KS-Varianz zu Pooled oder Total-Varianz der Datenverlust in der Berechnung etwas reduziert werden kann: die Pooled Varianz nutzt immerhin fast 80% der nötigen Klassen, die Totale Varianz bereits fast 82%. Alleinig die Ansätze, bei denen die Totale Varianz-Schätzung mit der Imputation ergänzt wird (Ansatz 8 und 10), ermöglichen eine 100% Wiederherstellung der Klassen, sodass alle Daten mit einbezogen werden können.

Betrachtet man die Stabilität des R_{max} , dann wird deutlich, dass die Werte über die Ansätze hinweg nur geringfügig vom klassischen KS-Ansatz (Ansatz 1) abweichen. Dabei sind die Abweichungen vom Referenzansatz bei den Ergebnissen der Pooled Varianz in der Tendenz etwas geringer als bei den Ansätzen mit der Totalen Varianz. Diese Beobachtung ist nachvollziehbar, da die Totale Varianz aufgrund ihrer Definition (Theorem 3) grösser ist als die Varianz des KS (S. 7, KS 5.3) oder der Pooled Varianz (Theorem 2). Dementsprechend ist auch ein höheres R_{max} zu erwarten. Abweichungen zwischen den R_{max} entstehen aber nicht nur aufgrund der unterschiedlichen Definition der verwendeten Varianzen, sondern auch durch die Unterschiede in der verwendeten Datengrundlage. Diese führt auch zu geringfügigen Unterschieden in den Schätzern A und B bei der R_{max} -Berechnung. Für letzteres lässt sich aber kein eindeutiger Trend aus den Analysen ableiten.

Bei der Beurteilung der Schätzunsicherheit des R_{max} liegen die Ansätze der Pooled Varianz und Totalen Varianz im klassischen (Ansätze 2 und 3) als auch erweiterten (Ansätze 5 und 6) Ansatz vergleichsweise nah beim Referenzmodell. Dies zeigt, dass die Wahl des Ansatzes nur einen vergleichsweise geringen Einfluss

auf die Schätzunsicherheit hat. Werden allerdings die Modelle deutlich komplexer (z.B. bei einer Ergänzung mit Imputation; Ansätze 7-10), sollte sich die Schätzunsicherheit in der Erwartung nach oben gehen. In diesem Beispiel ist dies allerdings nur marginal zu sehen. Die Ergebnisse auf den anderen Tests im Appendix D, zeigen diesen Effekt etwas deutlicher.

In Bemerkung 6(b) wurde für den Imputationsansatz I1 darauf hingewiesen, dass aufgrund der Definition des Ansatzes, «gesunde Klassen» bei der Imputation vernachlässigt werden müssen und dies eventuell zu höheren R_{max} führen könnte. Vergleicht man die Ergebnisse der Ansätze mit Imputation I1 (d.h., Ansätze 7 und 8) mit den Ansätzen ohne Imputation (d.h., Ansätze 5 und 6), dann bestätigt sich dies nicht eindeutig. Dass der Effekt vergleichsweise gering zu sein scheint, ist vermutlich darauf zurückzuführen, dass in vielen Beispielen der Anteil an «gesunden Klassen» in der Basisversicherung (d.h., $NBase_k > 0$ und $LBase_k = 0$) im Vergleich zu der Gesamtzahl an Klassen, sehr gering ist und damit nicht sehr stark ins Gewicht fällt.

Dieses Ungleichgewicht von «gesunden Klassen» in der Basisversicherung zu Gesamtzahl der Klassen hatten wir auch bereits in Bemerkung 7(b) diskutiert. Unsere Erwartung ist, dass die Imputation I2 nur bedingt gut funktioniert, wenn genau dieses Ungleichgewicht in den bestehenden Daten existiert. Die Ergebnisse (vergleicht man jeweils Ansätze 9 und 10 mit den Ansätzen 5 und 6) zeigen, dass es hier wirklich zu jeweils grösseren Abweichungen kommt. Die Imputation scheint hier entsprechend nicht das bestehende Ergebnis (ohne Imputation) zu bestätigen. Eine Analyse der Gütekriterien der Regressionen in I2 für KK-AJ, Modell M29 (siehe Tabelle 12) bestätigt diesen beobachteten Sachverhalt.

	Anzahl Klassen verfügbar				Anzahl Klassen geschätzt		R^2 adjusted Logistische Regression
	Total	Fitting	Fitting gesund	Fitting krank	Prediction gesund	Prediction krank	
Kriterien	$NMC_k > 0$	$NBase_k > 0$	$NBase_k > 0$ & $LBase_k = 0$	$NBase_k > 0$ & $LBase_k > 0$	$NBase_k > 0$ & $LBase_k = 0$	$NBase_k > 0$ & $LBase_k > 0$	
KK-AJ, M29	104	85	5	80	18	1	83%

Tabelle 12 Gütekriterien zur Logistischen Regression in I2 für KK-AJ, M29

Insgesamt sind 104 Klassen für die Berechnung des Kostennachweises vorhanden (d.h., $NMC_k > 0$), wobei bei 85 Klassen auch entsprechende Informationen in der Basisversicherung verfügbar sind (d.h., $NBase_k > 0$ und entspricht den Daten aus Ansatz 6, bei dem ohne Imputation der maximal möglich Datenumfang verwendet werden kann). Diese Klassen können im Fitting der Logistischen Regression verwendet werden, während für die fehlenden 19 (= 104 – 85) Klassen durch die Imputation Werte zu bestimmen sind. Von den 85 Klassen, die für das Fitting zur Verfügung stehen, sind allerdings nur 5 sogenannte «gesunde Klassen» in der Basisversicherung (d.h., mit $NBase_k > 0$ und $LBase_k = 0$), d.h., das Verhältnis von 6% (gesund) zu 94% (krank) in den Fitting-Daten steht in einem deutlichen Ungleichgewicht. Wird nun die Logistische Regression durchgeführt, scheint die Näherung der Regression mit einem Gütemass von 83% vergleichsweise gut zu sein. Die Prediction auf den zu schätzenden 19 Klassen zeigt allerdings ein Verhältnis von 95% (gesund) zu 5% (krank), was genau dem Gegenteil der Beobachtungen auf dem gesamten Datensatz entspricht. Die Logistische Regression erzeugt hier also unrealistische Resultate, die nur schwer zu plausibilisieren sind. Aufgrund der Komplexität des Imputationsansatzes I2 und den vergleichweisen unzureichenden Ergebnissen, sehen wir die Ansätze im Imputationsansatz I1 (d.h., 7 und 8) klar im Vorteil.

Die Ergebnisse zeigen, dass mittels der Totalen Varianz auch ohne Imputation (Ansatz 6) bereits ein Grossteil der Daten zur Berechnung zugelassen werden kann. Auch wenn sich die Schätzunsicherheit bei diesem Ansatz im Vergleich zum Referenzansatz nur geringfügig zu ändern scheint, muss festgehalten werden, dass ein Ansatz mit Totaler Varianz zu einem geringfügig grösseren R_{max} führt. Dies ist eine erwartete Konsequenz, da die Totale Varianz grundsätzlich eine grössere Abweichung (Variation innerhalb und zwischen den Klassen) misst als dies bei der Vorgehensweise im gegenwärtigen KS 5.3 der Fall ist (es wird nur die Variation innerhalb der Klassen gemessen). Sollen die Ansätze weiter verfeinert werden und eine Imputation mit eingeführt werden, so ist der Imputationsansatz I1 klar vorzuziehen.

Die Beobachtungen auf Basis der quantitativen Auswertung decken sich mit den Resultaten der qualitativen Auswertungen (siehe Abschnitt 5.1).

6. Analyse Pharmazeutische Kostengruppen (PCG)

6.1. Einführung der Kostengruppen

Gemäss Auftragsverständnis ist eine der entscheidenden Schwachstellen im aktuellen KS 5.3, dass sich der Kostennachweis nicht auf die neuesten Entwicklungen im Bereich Risikoausgleich abstützt. Hier wird vor allem darauf verwiesen, dass die Pharmazeutischen Kostengruppen (Pharmaceutical Cost Groups = PCG), die im Jahr 2020 im Risikoausgleich mit aufgenommen wurden um ambulanten Leistungen mehr Gewicht zu geben, aktuell bei den Kostennachweisen nicht berücksichtigt werden. Die Vermutung des BAG ist, dass dadurch die Bestandsstrukturangleichung zwischen Basisversicherten (freie Wahl des Leistungserbringers) und Modellversicherten beeinträchtigt ist, was wiederum die Aussagekraft des maximal zulässigen Modellrabatt mindern könnte.

Um diesen Punkt zu analysieren, haben wir die Berechnungen aus Kapitel 5 nochmals durchgeführt, diesmal allerdings mit einer zusätzlichen Differenzierungsvariable für die Klassen im bestehenden EF-Datensatz, d.h., den PCG. Da es doch eine erhebliche Anzahl an verschiedenen PCG-Indikatoren gibt (z.B. im EFIND Datensatz, siehe Abschnitt 5.2.3, sind dies insgesamt 35), muss eine Aggregation der PCG-Kategorien erfolgen. Würden alle PCG-Indikatoren einzeln in die Analysen (d.h., der Erstellung der MF-MC Daten) einfließen, würde eine unüberschaubare Menge an Klassen entstehen, bei denen viele Klassen gar keine bis nur extrem geringe Beobachtungen beinhalten würden. So würde sich die Schwachstelle (2) weiter ausprägen.

Deshalb haben wir in einem ersten Ansatz die naheliegendste Segmentierung vorgenommen: Gruppierung #1 in drei PCG-Klassen mit «Keinem PCG-Flag», «Genau einem PCG-Flag» und «Mehr als einem PCG-Flag». Um die Sensitivität der Variablen PCG zu testen, haben wir eine alternative Segmentierung (Gruppierung #2) in den Berechnungen eingefügt:

- (a) «Keinem PCG-Flag»
- (b) «Genau einem PCG-Flag in PCG_NIE» (Nierenerkrankung)
- (c) «Genau einem PCG-Flag in PCG_PAH» (Pulmonale (arterielle) Hypertonie)
- (d) «Genau einem PCG-Flag in PCG_KRK» (Krebs komplex)
- (e) «Genau einem PCG-Flag in einer der anderen PCG als in (b), (c) oder (d)», und
- (f) «Mehr als einem PCG-Flag»

Bei den Segmenten (b) – (d) handelt es sich um die in Bezug auf Gesamtkosten teuersten Erkrankungen in der Stichprobe. Diese Analyse wurde durch eine separate Lineare Regression auf dem EFIND-Datensatz bestätigt. Eine weitere Verifizierung dieser Beobachtung konnten durch die Statistik zum Risikoausgleich (2020) erhalten werden. Bei den Arzneimitteldaten sind die monatlichen Zuschläge auf genau diesen drei PCG erhöht. Die Segmentierungen wurden mit dem PuS diskutiert und zur Kenntnis genommen.

6.2. Bewertung der Ansätze bei Einbezug der Kostengruppen (PCG)

Aus qualitativer Optik lässt sich festhalten, dass sich – durch die Einführung der PCG als zusätzliches Segmentierungsmerkmal im MF-MC Datensatz – die Anzahl der zur Verfügung stehenden Klassen deutlich erhöhen wird. Bei einer Segmentierung «Keinen PCG-Flag», «Genau einem PCG-Flag» und «Mehr als einem PCG-Flag» erhöht sich die Anzahl der zur Verfügung stehenden Klassen (falls alle Klassen besetzt wären) bereits um den Faktor drei. Dies führt aber auch dazu, dass die einzelnen Kollektive pro Klasse (unabhängig ob Basisversicherten- oder Modellversicherten-Kollektiv) kleiner werden. In diesem Fall wird

Schwachstelle (2) verschärft. Ein entsprechendes Beibehalten der aktuell gültigen Berechnung gemäss KS 5.3, alleinig durch die Ergänzung der PCG, scheint daher mittel- bis langfristig nicht zielführend.

Im Rahmen unserer Analysen haben wir deshalb die Berechnungen aus Kapitel 5 auf dem Imputationsmodell mit Totaler Varianz (Ansatz 8) als auch für das originale Berechnungskonstrukt des KS 5.3 (Referenzmodell/Ansatz 1) nochmals nach Einführung und Segmentierung der PCGs durchgeführt.

Die folgende Tabelle zeigt die Ergebnisse der Analyse für die beiden alternativen PCG-Segmentierungen im Vergleich zum Fall ohne PCG (also Ergebnisse Abschnitt 5.3), wieder beispielhaft für Krankenversicherer KK-AJ, im Modell M29:

	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)/(8)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
KS Klassisch – Ansatz (1)								
Ohne PCG	70	104	67.3%	988.82		78'669.64	280.48	28.4%
Mit PCG-Gruppierung #1	85	169	50.3%	721.33	-27.1%	67'884.06	260.55	36.1%
Mit PCG-Gruppierung #2	83	172	48.3%	646.57	-34.6%	67'585.26	259.97	40.2%
Total erweitert I1 – Ansatz (8)								
Ohne PCG; Ansatz (8)	104	104	100.0%	1'024.73		84'982.95	291.52	28.4%
Mit PCG-Gruppierung #1	169	169	100.0%	937.05	-8.6%	84'409.35	290.53	31.0%
Mit PCG-Gruppierung #2	172	172	100.0%	814.60	-20.5%	68'036.07	260.84	32.0%

Tabelle 13 Quantitative Bewertung der PCG anhand von KK-AJ für Modell M29

Wie erwartet zeigt sich, dass die Einführung der PCG als zusätzliches Klassenkriterium die zur Verwendung stehenden Daten im Referenzmodell weiter einschränkt (von bisher 67% verwendeten Klassen ohne PCG, auf 50%/48% mit PCG). Der bestehende Ansatz im KS 5.3 inklusive PCG stellt also mittel- bis langfristig keine geeignete Option dar. Einen eindeutigen Unterschied bezüglich der Art der PCG-Segmentierung scheint es nicht zu geben.

Was allerdings zusätzlich auffällt ist, dass die Höhe des Maximalrabatts R_{max} bei der Einführung von PCG sich deutlich verändert. Die Höhe der Veränderungen scheinen allerdings unabhängig vom verwendeten Ansatz zu sein (siehe dazu auch die Ergebnisse in Appendix D) und erscheinen eher zufällig. Auch die Betrachtung der Variationskoeffizienten zeigt keine eindeutigen Trends in der relativen Unschärfe. Nach diesen Kriterien scheinen die beiden Ansätze kompatibel zu sein. Allerdings ist festzuhalten, dass beim Ansatz 8 wieder alle zur Verfügung stehenden Klassen in die Berechnung des Kostennachweises einbezogen werden können.

Appendix A. - *Abkürzungen*

Relevante Abkürzungen:

ANOVA	Analysis of Variance
BAG	Bundesamt für Gesundheit
KN	Kostennachweis
KS	Kreisschreiben (hauptsächlich Bezug hier auf KS 5.1 und 5.3 betreffend die obligatorische Krankenversicherung)
KVAG	Krankenversicherungsaufsichtsgesetz (Bundesgesetz betreffend die Aufsicht über die soziale Krankenversicherung)
KVAV	Krankenversicherungsaufsichtsverordnung (Verordnung betreffend die Aufsicht über die soziale Krankenversicherung)
KVG	Krankenversicherungsgesetz (Bundesgesetz über die Krankenversicherung)
KVV	Krankenversicherungsverordnung (Verordnung über die Krankenversicherung)
OKP	Obligatorische Krankenpflegeversicherung
PCG	Pharmazeutischen Kostengruppen
PuS	Abteilung Prämien und Solvenzaufsicht innerhalb des BAG
PwC	PricewaterhouseCoopers AG, Zürich

Appendix B. - *BAG-Daten und Lieferergebnisse*

BAG Pseudo-Daten:

- 1111-EF_MC_2021_d.xlsx
- BAG_Zihler_Modellrabatte_Testdata_2020_20221018_downloaded_Versand_.csv
- BAG_Zihler_Modellrabatte_Testdata_2021_20221018_downloaded_Versand_.csv

Lieferergebnisse:

- Präsentationen zu den wöchentlichen Meetings:
 - PwC_Resultate_130323.pdf (versendet am 14.03.2023)
 - PwC_Franchise_Comparison_090323.pdf (versendet am 09.03.2023)
 - PwC_Franchise_Comparison_070323.pdf (versendet am 08.03.2023)
 - PwC_Resultate_210223.pdf (versendet am 20.02.2023)
 - PwC support BAG 20230110.pdf (gezeigt am 24.01.2023; versendet am 2. Mai 2023)
 - PwC support BAG 20221214.pdf (versendet am 15.12.2022)
 - PwC support BAG 20221207.pdf (versendet am 09.12.2022)
 - PwC support BAG 20221130.pdf (versendet am 30.11.2022)
 - PwC support BAG 20221123_RP.pdf (versendet am 23.11.2022)
 - TotalVarianz_151122.pdf (versendet am 15.11.2022)
 - PwC support BAG 20221116.pdf (gezeigt am 16.11.2022)
 - PwC support BAG 20221109.pdf (versendet am 10.11.2022)
 - PwC support BAG 20221102_v1.pdf (versendet am 1.11.2022)
- R-Codes:
 - auswertung_efmc_UTF8_Final_3.R (14.04.2023)
 - user_input_Final.R (Stand: 17.03.2023)
 - 1_mc_daten_lesen_UTF8_Final_2.R (Stand: 24.03.2023)
 - 2_rueckmeldung_efmc_UTF8_Final_3.R (Stand: 19.04.2023)
 - 4_run_mc_check_UTF8_Final_3.R (Stand: 19.04.2023)
 - 5_plot_and_table_Final.R (Stand: 21.03.2023)
 - 0_EFIND_Vorbereitung_v7.R (wurde dem BAG für die Tests in diesem Bricht zur Verfügung gestellt; dies ist aber kein Bestandteil der finalen Lieferergebnisse)

Appendix C. - *Beweise*

Beweis Theorem 1:

Der Beweis beruht auf der Tatsache, dass bei der Betrachtung von Mittelwerten oft der Kreuzterm verschwindet.

$$\begin{aligned}\sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y})^2 &= \sum_{i=1}^K \sum_{j=1}^{n_i} ((y_{i,j} - \bar{y}_i) + (\bar{y}_i - \bar{y}))^2 \\&= \sum_{i=1}^K \sum_{j=1}^{n_i} ((y_{i,j} - \bar{y}_i)^2 + 2 \cdot (y_{i,j} - \bar{y}_i) \cdot (\bar{y}_i - \bar{y}) + (\bar{y}_i - \bar{y})^2) \\&= \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i)^2 + 2 \cdot \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i) \cdot (\bar{y}_i - \bar{y}) + \sum_{i=1}^K n_i \cdot (\bar{y}_i - \bar{y})^2 \\&= \sum_{i=1}^K n_i \cdot (\bar{y}_i - \bar{y})^2 + \sum_{i=1}^K \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i)^2 + 2 \cdot \sum_{i=1}^K (\bar{y}_i - \bar{y}) \sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i)\end{aligned}$$

mit

$$\sum_{j=1}^{n_i} (y_{i,j} - \bar{y}_i) = \sum_{j=1}^{n_i} y_{i,j} - \sum_{j=1}^{n_i} \bar{y}_i = \sum_{j=1}^{n_i} y_{i,j} - n_i \cdot \bar{y}_i = 0$$

folgt dann das Ergebnis. □

Beweis Theorem 2:

Teil 1 – Varianz eines Schätzers $\bar{X}$:

Nach KS 5.3 muss die Varianz eines Schätzers A bzw. B bestimmt werden, bei dem es sich in beiden Fällen um einen Schätzer eines Durchschnitts handelt. Nehmen wir deshalb klassisch an, dass wir die Varianz eines arithmetischen Mittels $\bar{X} = \frac{1}{N} \sum_{i=1}^N X_i$ bestimmen wollen. Unter der Annahme, dass die X_i unabhängig und identisch verteilt sind, folgt für $Var(\bar{X})$:

$$Var(\bar{X}) = Var\left(\frac{1}{N} \sum_{i=1}^N X_i\right) = \frac{1}{N^2} \sum_{i=1}^N Var(X_i) = \frac{1}{N^2} \cdot N \cdot Var(X_1) = \frac{1}{N} \cdot Var(X_1).$$

Die $Var(X_1)$ kann dann beispielsweise über die Stichprobenvarianz der X_i berechnet werden.

Teil 2 – $Var(A)$:

Nach Theorem 1, der Definition in Tabelle 4, als auch Teil 1 dieses Beweises folgt

$$Var(A) = \frac{1}{NMC} \cdot \frac{1}{NMC - K} \cdot SSE_A$$

mit

$$SSE_A = \sum_{k=1}^K \sum_{i=1}^{NMC_k} \left(LMC_{k,i} - \frac{LMC_k}{NMC_k} \right)^2 .$$

Die Variablen K , NMC , NMC_k , LMC_k als auch $LMC_{k,i}$ sind analog definiert wie in Abschnitt 3.1.

Des Weiteren kann SSE_A noch wie folgt umformuliert werden:

$$\begin{aligned} SSE_A &= \sum_{k=1}^K \sum_{i=1}^{NMC_k} \left(LMC_{k,i} - \frac{LMC_k}{NMC_k} \right)^2 = \sum_{k=1}^K \sum_{i=1}^{NMC_k} \left(LMC_{k,i}^2 - 2 \cdot LMC_{k,i} \frac{LMC_k}{NMC_k} + \frac{LMC_k^2}{NMC_k^2} \right) = \\ &= \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}^2 - 2 \sum_{k=1}^K \frac{LMC_k}{NMC_k} \sum_{i=1}^{NMC_k} LMC_{k,i} + \sum_{k=1}^K NMC_k \cdot \frac{LMC_k^2}{NMC_k^2} . \end{aligned}$$

Mit den Definitionen aus dem KS 5.3 $\sum_{i=1}^{NMC_k} LMC_{k,i}^2 = QMC_k$ und $\sum_{i=1}^{NMC_k} LMC_{k,i} = LMC_k$ folgt dann

$$SSE_A = \sum_{k=1}^K QMC_k - 2 \sum_{k=1}^K \frac{LMC_k^2}{NMC_k} + \sum_{k=1}^K \frac{LMC_k^2}{NMC_k} = \sum_{k=1}^K \left(QMC_k - \frac{LMC_k^2}{NMC_k} \right) .$$

Teil 3 – Var(B):

Für diesen Teil nehmen wir zuerst noch eine Umformulierung des Schätzers B in der folgenden Form vor:

$$\begin{aligned} B &= \frac{1}{NMC} \sum_{k=1}^K NMC_k \frac{LBase_k}{NBase_k} = \frac{NBase}{NBase} \cdot \frac{1}{NMC} \sum_{k=1}^K \frac{NMC_k}{NBase_k} \sum_{i=1}^{NBase_k} LBase_{k,i} \\ &= \frac{1}{NBase} \sum_{k=1}^K \frac{NBase}{NMC} \cdot \frac{NMC_k}{NBase_k} \sum_{i=1}^{NBase_k} LBase_{k,i} = \frac{1}{NBase} \sum_{k=1}^K \sum_{i=1}^{NBase_k} LBase_{k,i} \cdot c_k \\ &= \frac{1}{NBase} \sum_{k=1}^K c_k \cdot LBase_k , \end{aligned}$$

mit $c_k := b \cdot a_k$, $b := \frac{NBase}{NMC}$ und $a_k := \frac{NMC_k}{NBase_k}$. Die entsprechenden Parameter K , NMC , $NBase$, NMC_k , $NBase_k$, $LBase_{k,i}$ als auch $QBase_k$ sind auch hier analog zu Abschnitt 3.1 definiert.

Die Berechnung der $Var(B)$ folgt dann analog zu Teil (2):

$$Var(B) = \frac{1}{NBase} \cdot \frac{1}{NBase - K} \cdot SSE_B$$

mit

$$SSE_B = \sum_{k=1}^K \sum_{i=1}^{NBase_k} c_k^2 \cdot \left(LBase_{k,i} - \frac{LBase_k}{NBase_k} \right)^2 ,$$

wobei die entsprechenden Beobachtungen $y_{i,j}$ in Theorem 1 durch $c_k \cdot LBase_{k,i}$ ersetzt werden.

Analog zur Umformung unter Teil (2) lässt sich dann SSE_B umformen in

$$SSE_B = \sum_{k=1}^K c_k^2 \sum_{i=1}^{NBase_k} \left(LBase_{k,i} - \frac{LBase_k}{NBase_k} \right)^2 = \sum_{k=1}^K c_k^2 \cdot \left(QBase_k - \frac{LBase_k^2}{NBase_k} \right) .$$

□

Beweis Theorem 3:

Im Folgenden sind die verwendeten Parameter und Variablen (d.h., K , NMC , etc.) analog zu den Definitionen in Abschnitt 3.1 bzw. wie für den Beweis zum Theorem 2 definiert.

Berechnungen erfolgen analog wie in Theorem 2, allerdings unter Berücksichtigung der SST_A und SST_B und den folgenden Umformungen:

$$\begin{aligned} SST_A &= \sum_{k=1}^K \sum_{i=1}^{NMC_k} (LMC_{k,i} - A)^2 = \sum_{k=1}^K \sum_{i=1}^{NMC_k} (LMC_{k,i}^2 - 2 \cdot LMC_{k,i} \cdot A + A^2) = \\ &= \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}^2 - 2 \cdot A \cdot \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i} + NMC \cdot A^2 = \\ &= \sum_{k=1}^K QMC_k - 2 \cdot A \cdot NMC \cdot \underbrace{\frac{1}{NMC} \cdot \sum_{k=1}^K \sum_{i=1}^{NMC_k} LMC_{k,i}}_{=A} + NMC \cdot A^2 \\ &= \sum_{k=1}^K QMC_k - 2 \cdot NMC \cdot A^2 + NMC \cdot A^2 \\ &= \sum_{k=1}^K QMC_k - NMC \cdot A^2 = \sum_{k=1}^K QMC_k - NMC \cdot \frac{(\sum_{k=1}^K LMC_k)^2}{NMC^2} = \sum_{k=1}^K QMC_k - \frac{(\sum_{k=1}^K LMC_k)^2}{NMC} \end{aligned}$$

und

$$\begin{aligned} SST_B &= \sum_{k=1}^K \sum_{i=1}^{NBase_k} (LBase_{k,i} \cdot c_k - B)^2 = \sum_{k=1}^K \sum_{i=1}^{NBase_k} (LBase_{k,i}^2 \cdot c_k^2 - 2 \cdot LBase_{k,i} \cdot c_k \cdot B + B^2) \\ &= \sum_{k=1}^K c_k^2 \cdot QBase_k - 2 \cdot B \cdot \sum_{k=1}^K c_k \sum_{i=1}^{NBase_k} LBase_{k,i} + NBase \cdot B^2 \\ &= \sum_{k=1}^K c_k^2 \cdot QBase_k - 2 \cdot B \sum_{k=1}^K c_k \cdot LBase_k + NBase \cdot B^2 = \\ &= \sum_{k=1}^K c_k^2 \cdot QBase_k - 2 \cdot B \cdot NBase \cdot \underbrace{\frac{1}{NBase} \sum_{k=1}^K c_k \cdot LBase_k}_{=B} + NBase \cdot B^2 \end{aligned}$$

$$\begin{aligned}
&= \sum_{k=1}^K c_k^2 \cdot QBase_k - NBase \cdot B^2 = \sum_{k=1}^K c_k^2 \cdot QBase_k - NBase \cdot \frac{(\sum_{k=1}^K c_k \cdot LBase_k)^2}{NBase^2} = \\
&= \sum_{k=1}^K c_k^2 \cdot QBase_k - \frac{(\sum_{k=1}^K c_k \cdot LBase_k)^2}{NBase} .
\end{aligned}$$

□

Appendix D. - Weitere Analyseergebnisse

Ergänzende Analysen zu Abschnitt 5.3 für Versicherer KK-AE, KK-AJ und KK-AD mit den jeweiligen Modellen M146 für KK-AE, M30 für KK-AJ und M117 und M150 für KK-AD:

Ansatz	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	total	in %					
(1) KS Klassisch	187	353	53.0%	429.85		21'933.74	148.10	34.5%
(2) Pooled Klassisch	187	353	53.0%	421.25	-2.0%	22'336.65	149.45	35.5%
(3) Total Klassisch	187	353	53.0%	454.07	5.6%	23'471.30	153.20	33.7%
(4) KS erweitert	195	353	55.2%	433.52	0.9%	21'798.74	147.64	34.1%
(5) Pooled erweitert	268	353	75.9%	487.50	13.4%	26'975.74	164.24	33.7%
(6) Total erweitert	295	353	83.6%	537.01	24.9%	29'574.78	171.97	32.0%
(7) Pooled erweitert I1	308	353	87.3%	477.98	11.2%	27'203.95	164.94	34.5%
(8) Total erweitert I1	353	353	100.0%	514.10	19.6%	29'964.04	173.10	33.7%
(9) Pooled erweitert I2	308	353	87.3%	448.47	4.3%	27'620.89	166.20	37.1%
(10) Total erweitert I2	353	353	100.0%	483.00	12.4%	29'009.38	170.32	35.3%

Tabelle 14 Quantitative Bewertung KK-AE für Modell M146

Ansatz	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	total	in %					
(1) KS Klassisch	55	72	76.4%	946.33		51'633.48	227.23	24.0%
(2) Pooled Klassisch	55	72	76.4%	827.39	-12.6%	42'163.31	205.34	24.8%
(3) Total Klassisch	55	72	76.4%	855.61	-9.6%	43'072.37	207.54	24.3%
(4) KS erweitert	55	72	76.4%	946.33	0.0%	51'633.48	227.23	24.0%
(5) Pooled erweitert	63	72	87.5%	940.28	-0.6%	48'859.43	221.04	23.5%
(6) Total erweitert	64	72	88.9%	991.37	4.8%	51'567.26	227.08	22.9%
(7) Pooled erweitert I1	70	72	97.2%	953.56	0.8%	52'508.84	229.15	24.0%
(8) Total erweitert I1	72	72	100.0%	1'011.67	6.9%	57'248.19	239.27	23.7%
(9) Pooled erweitert I2	70	72	97.2%	910.12	-3.8%	47'772.80	218.57	24.0%
(10) Total erweitert I2	72	72	100.0%	961.09	1.6%	50'104.25	223.84	23.3%

Tabelle 15 Quantitative Bewertung KK-AJ für Modell M30

Ansatz	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	total	in %					
(1) KS Klassisch	1'987	3'120	63.7%	386.46		1'008.95	31.76	8.2%
(2) Pooled Klassisch	1'987	3'120	63.7%	367.11	-5.0%	954.95	30.90	8.4%
(3) Total Klassisch	1'987	3'120	63.7%	373.44	-3.4%	959.23	30.97	8.3%
(4) KS erweitert	2'061	3'120	66.1%	386.08	-0.1%	1'021.58	31.96	8.3%
(5) Pooled erweitert	2'719	3'120	87.1%	388.42	0.5%	1'055.92	32.49	8.4%
(6) Total erweitert	2'882	3'120	92.4%	399.33	3.3%	1'136.79	33.72	8.4%
(7) Pooled erweitert I1	2'867	3'120	91.9%	387.14	0.2%	1'068.78	32.69	8.4%
(8) Total erweitert I1	3'120	3'120	100.0%	396.79	2.7%	1'161.67	34.08	8.6%
(9) Pooled erweitert I2	2'867	3'120	91.9%	380.21	-1.6%	1'061.79	32.59	8.6%
(10) Total erweitert I2	3'120	3'120	100.0%	386.38	0.0%	1'310.94	36.21	9.4%

Tabelle 16 Quantitative Bewertung KK-AD für Modell M117

Ansatz	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	total	in %					
(1) KS Klassisch	3'234	5'111	63.3%	331.28		1'098.16	33.14	10.0%
(2) Pooled Klassisch	3'234	5'111	63.3%	327.04	-1.3%	1'101.24	33.18	10.1%
(3) Total Klassisch	3'234	5'111	63.3%	337.92	2.0%	1'102.00	33.20	9.8%
(4) KS erweitert	3'397	5'111	66.5%	338.76	2.3%	1'119.26	33.46	9.9%
(5) Pooled erweitert	3'981	5'111	77.9%	348.23	5.1%	1'175.30	34.28	9.8%
(6) Total erweitert	4'384	5'111	85.8%	383.86	15.9%	1'366.89	36.97	9.6%
(7) Pooled erweitert I1	4'431	5'111	86.7%	345.75	4.4%	1'187.22	34.46	10.0%
(8) Total erweitert I1	5'111	5'111	100.0%	365.83	10.4%	1'391.24	37.30	10.2%
(9) Pooled erweitert I2	4'431	5'111	86.7%	331.03	-0.1%	1'467.18	38.30	11.6%
(10) Total erweitert I2	5'111	5'111	100.0%	344.82	4.1%	1'963.07	44.31	12.8%

Tabelle 17 Quantitative Bewertung KK-AD für Modell M150

	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)/(8)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
KS Klassisch – Ansatz (1)								
Ohne PCG	187	353	53.0%	429.85		21'933.74	148.10	34.5%
Mit PCG-Gruppierung #1	251	566	44.3%	499.56	16.2%	19'598.77	140.00	28.0%
Mit PCG-Gruppierung #2	250	573	43.6%	458.07	6.6%	20'002.76	141.43	30.9%
Total erweitert I1 – Ansatz (8)								
Ohne PCG; Ansatz (8)	353	353	100.0%	514.10		29'964.04	173.10	33.7%
Mit PCG-Gruppierung #1	566	566	100.0%	608.27	18.3%	36'804.88	191.85	31.5%
Mit PCG-Gruppierung #2	573	573	100.0%	567.50	10.4%	37'319.74	193.18	34.0%

Tabelle 18 Quantitative Bewertung der PCG anhand von KK-AE für Modell M146

	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)/(8)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
KS Klassisch – Ansatz (1)								
Ohne PCG	55	72	76.4%	946.33		51'633.48	227.23	24.0%
Mit PCG-Gruppierung #1	59	105	56.2%	524.43	-44.6%	49'362.41	222.18	42.4%
Mit PCG-Gruppierung #2	59	107	55.1%	593.53	-37.3%	49'151.03	221.70	37.4%
Total erweitert I1 – Ansatz (8)								
Ohne PCG; Ansatz (8)	72	72	100.0%	1'011.67		57'248.19	239.27	23.7%
Mit PCG-Gruppierung #1	105	105	100.0%	631.57	-37.6%	44'948.77	212.01	33.6%
Mit PCG-Gruppierung #2	107	107	100.0%	611.69	-39.5%	44'658.85	211.33	34.5%

Tabelle 19 Quantitative Bewertung der PCG anhand von KK-AJ für Modell M30

	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)/(8)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
KS Klassisch – Ansatz (1)								
Ohne PCG	1'987	3'120	63.7%	386.46		1'008.95	31.76	8.2%
Mit PCG-Gruppierung #1	2'419	4'601	52.6%	201.22	-47.9%	749.95	27.39	13.6%
Mit PCG-Gruppierung #2	2'416	4'637	52.1%	195.76	-49.3%	701.31	26.48	13.5%
Total erweitert I1 – Ansatz (8)								
Ohne PCG; Ansatz (8)	3'120	3'120	100.0%	396.79		1'161.67	34.08	8.6%
Mit PCG-Gruppierung #1	4'601	4'601	100.0%	213.99	-46.1%	1'165.68	34.14	16.0%
Mit PCG-Gruppierung #2	4'637	4'637	100.0%	184.57	-53.5%	1'036.20	32.19	17.4%

Tabelle 20 Quantitative Bewertung der PCG anhand von KK-AD für Modell M117

	Anzahl Klassen			R_{max}	Abweichung R_{max} von (1)/(8)	$Var(R_{max})$	$\sqrt{Var(R_{max})}$	$\frac{\sqrt{Var(R_{max})}}{R_{max}}$
	verwendete	Total	in %					
KS Klassisch – Ansatz (1)								
Ohne PCG	3'234	5'111	63.3%	331.28		1'098.16	33.14	10.0%
Mit PCG-Gruppierung #1	5'003	9'421	53.1%	270.59	-18.3%	1'026.67	32.04	11.8%
Mit PCG-Gruppierung #2	4'988	9'642	51.7%	262.02	-20.9%	924.74	30.41	11.6%
Total erweitert I1 – Ansatz (8)								
Ohne PCG; Ansatz (8)	5'111	5'111	100.0%	365.83		1'391.24	37.30	10.2%
Mit PCG-Gruppierung #1	9'421	9'421	100.0%	238.03	-34.9%	1'179.04	34.34	14.4%
Mit PCG-Gruppierung #2	9'642	9'642	100.0%	164.67	-55.0%	1'181.94	34.38	20.9%

Tabelle 21 Quantitative Bewertung der PCG anhand von KK-AD für Modell M150

Referenzen

- Brownlee, J. (2020): «A Gentle Introduction to Threshold-Moving for Imbalanced Classification»; <https://machinelearningmastery.com/threshold-moving-for-imbalanced-classification/>; letzter Aufruf 21.04.2023.
- Casella, G. und Berger, R. L. (2002): «Statistical Inference, 2nd edition», Duxbury Advanced Series.
- Efron, B. und Tibshirani, R. J. (1994): «An Introduction to the Bootstrap»; Chapman & Hall/CRC Monographs on Statistics & Applied Probability; Kapitel 9, «Regression Models».
- Kreisschreiben 5.3 (2019): «Besondere versicherungsform mit eingeschränkter Wahl der Leistungserbringer: Nachweis Kostenunterschiede»; https://www.bag.admin.ch/dam/bag/de/dokumente/kuv-aufsicht/rakv2/ks-05-3-bes-form-wahl-le1.pdf.download.pdf/Kreisschreiben%205.3_d.pdf; letzter Aufruf 28.03.2023.
- Statistik zum Risikoausgleich (2020): «Arzneimitteldaten Eingruppierung: 2020»; <https://www.kvg.org/versicherer/risikoausgleich/risikoausgleich-pcg/>; letzter Aufruf 28.03.2023.
- Wollschläger, D. (2017): «Grundlagen der Datenanalyse mit R – Eine anwendungsorientierte Einführung»; Publisher: Springer Spektrum Berlin, Heidelberg; Series Titel: Statistik und ihre Anwendungen; DOI: <https://doi.org/10.1007/978-3-662-53670-4>.