

Projekt IGLAD: Entstehung und Zukunft einer internationalen In-Depth Unfalldatenbank

Abstract

Während Unfallstatistiken auf nationaler Ebene von vielen Ländern erhoben werden um Trends in der Entwicklung der Anzahl von Straßenverkehrsunfällen aufzuzeigen, gibt es einen großen Bedarf an Unfalldaten mit größerer Informationstiefe, so genannten In-Depth Daten. Aus diesem Grund entstanden in den letzten Jahren verschiedene Unfalldatenprojekte weltweit. Vergleichende Auswertungen dieser Daten waren bisher jedoch nicht möglich, da sich verschiedene Standards etabliert hatten. Das Projekt IGLAD (Initiative for the Global Harmonization of Accident Data) schaffte in den letzten zwei Jahren die Voraussetzungen für einen standardisierten Datensatz als gemeinsamer Nenner verschiedener Unfalldatenbanken in Europa, USA und Asien. Neben der Definition eines geeigneten Datenschemas und der Durchführung einer Pilotstudie wurde ein Geschäftsmodell entwickelt, das die Grundlage für ein sich selbst tragendes Projekt geschaffen hat und ein tragfähiges Konsortium zur weiteren Pflege der Daten wurde gegründet. Dieser Beitrag berichtet über den Verlauf und Status des Projektes, die sich stellenden Herausforderungen und die nächsten Schritte zur Schaffung einer harmonisierten Datenbank.

Motivation

Die Anzahl und Entwicklung der Verkehrstoten weltweit kann nur grob geschätzt werden und betrug für 2010 laut WHO 1.240.000. Der weitere Verlauf dieser Zahlen bis 2020 und darüber hinaus hängt von vielen Faktoren ab. Die bisherige Entwicklung in „High-income countries“ lässt darauf schließen, dass insbesondere durch Arbeit an Infrastruktur, Fahrzeugtechnik und Gesetzgebung ein wesentlicher Beitrag zur Verbesserung der Situation geleistet werden kann. Auch erste „In-Depth“ Unfalldaten aus Ländern wie Indien und China bestätigen dies. Abbildung 1 zeigt die Entwicklung der Verkehrstoten von 1990 ab bis 2010 als prognostische Weiterführung nach dem TFEC Modell [1]. Das Traffic Fatalities and Economic Growth (TFEC) Projekt der World Bank entwickelte dieses Modell auf Basis von Verkehrs-, Bevölkerungs- und ökonomischen Daten [2].

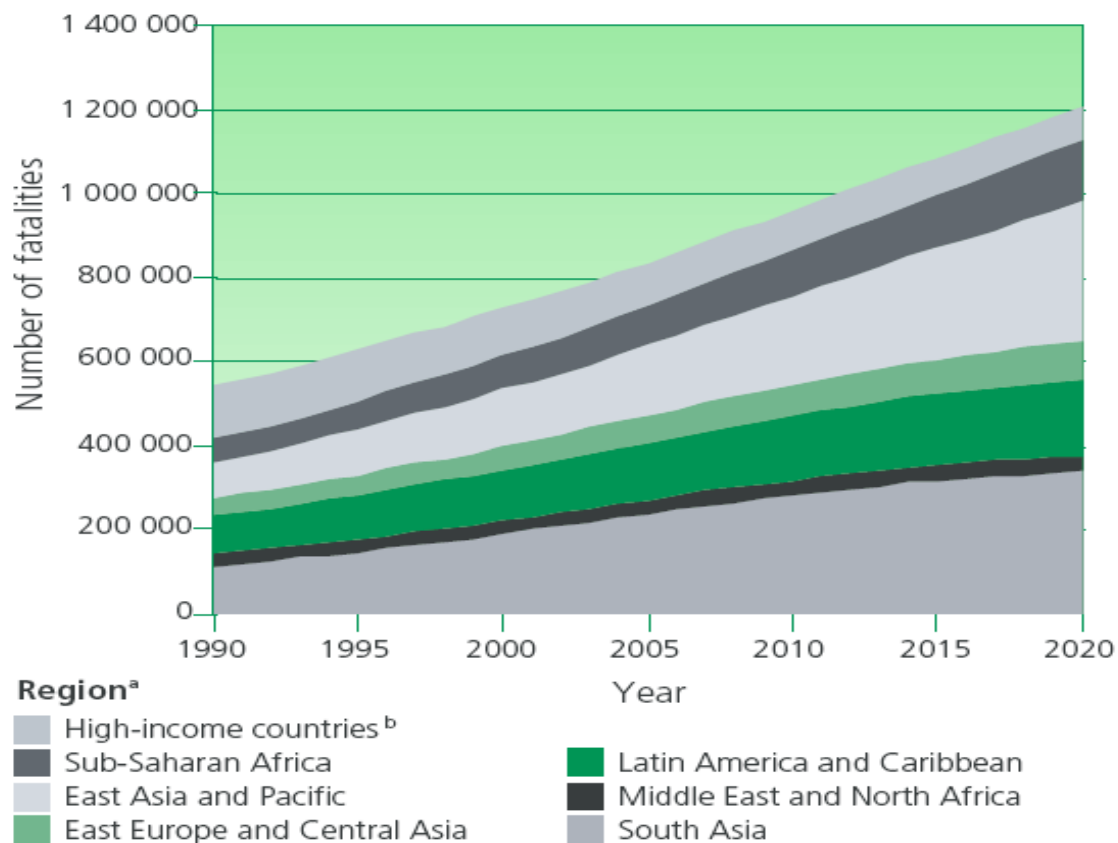


Abbildung 1: Entwicklung der Verkehrstoten weltweit (Quelle: WHO Bericht 2004).

Schon bei diesen grundlegenden Betrachtungen wird ersichtlich, dass es ein erhebliches Defizit an Daten gibt, um eine differenzierte und umfassende Einschätzung der Situation zu ermöglichen. Auf der anderen Seite hat sich herausgestellt, dass eine gute und detailreiche Datenbasis notwendig ist, um gezielt und effektiv die Verkehrsunfallsituation zu verbessern. Insbesondere bei der Verbesserung der Fahrzeugsicherheit sind Realunfalldaten eine wesentliche Grundlage. Aus diesem Grund haben sich in „High-Income countries“ schon seit einigen Jahren Projekte etabliert, die eine umfassende Erhebung von Unfalldaten durchführen.

Die Europäische Union hat sich mit dem Road Safety Action Programme [3] die Halbierung der Verkehrstoten jeweils in der ersten und zweiten Dekade dieses Jahrtausends zum Ziel gesetzt. Eine erfolgreiche Umsetzung wird zunehmend schwieriger, da die Wirkung einzelner Maßnahmen geringer wird, der Aufwand hingegen steigt. Deshalb ist es umso wichtiger, das Verkehrsunfallgeschehen genau zu analysieren und dafür eine geeignete europaweite Datengrundlage zu schaffen.

Deutlich wird dies bereits in dem 2010 eingeführten Euro NCAP Advanced reward [4] für neue Entwicklungen von Sicherheitssystemen in Pkw. Dieses Programm möchte

vielversprechende Technologien honorieren, die eine weitere Verringerung der Anzahl der Straßenverkehrsoffer ermöglichen. Eine Bewertung der Systeme erfolgt aufgrund von Effektivitätsanalysen, die auf Basis von Unfalldaten erstellt werden. Hier zeigt sich schnell, dass eine genaue Hochrechnung einer solchen Effektivität auf Europa nur mit umfassenden Daten aus verschiedenen europäischen Ländern möglich ist, da sich das Verkehrsunfallgeschehen in den einzelnen Ländern erheblich unterscheiden kann.

In-depth Unfallerhebungen konnten sich schon in einigen Ländern weltweit über einen längeren Zeitraum etablieren. Vor allem in neuen Märkten entstehen zurzeit weitere Erhebungen. Die einzelnen Erhebungsprojekte unterscheiden sich stark in Anzahl der aufgenommenen Fälle pro Jahr, Erhebungsumfang (Anzahl der Faktoren), Stichprobencharakteristik, Organisation und Finanzierung. Dies erschwert den Vergleich von Daten aus verschiedenen Ländern erheblich und z.B. Analysen, die Aussagen für eine größere Region wie Europa treffen sollen, sind zum einen recht ungenau und erfordern zum anderen ein umfassendes Wissen mehrerer Datenbanken und deren Analyse.

Die nationalen Statistiken eines Landes erfüllen in diesem Kontext zwei wichtige Funktionen: Sie bilden ein verbindendes Element zwischen den „In-Depth“ Stichproben des jeweiligen Landes und sie ermöglichen als Randverteilung eine Hochrechnung der „In-Depth“ Daten auf die jeweilige Nation. Eine Grundvoraussetzung für die Hochrechnung auf nationale Ebene ist, dass die Variablen der nationalen Daten auch in der „In-Depth“ Stichprobe vorhanden sein müssen. Vereinheitlichte und harmonisierte nationale Daten ermöglichen dann eine Hochrechnung auf multinationale Ebenen.

Der Harmonisierungsprozess bei nationalen Daten ist glücklicherweise schon recht fortgeschritten. Projekte wie IRTAD [5] und CARE [6] waren hier bereits sehr erfolgreich und liefern elementare Größen wie Anzahl der Getöteten, Anzahl der Unfälle mit Getöteten, Art der Verkehrsbeteiligung für unterschiedliche Länder in standardisierter Form, so dass z.B. auch die Definition von „getötet im Straßenverkehr“ vereinheitlicht ist. Durch Ergänzung von nationalen Daten mit harmonisierten „In-Depth“ Daten ist eine Erweiterung des Analysefeldes möglich. Harmonisierte „In-Depth“ Daten und nationale Daten zusammen ermöglichen einen Vergleich verschiedener Länder und eine Identifikation von Schwerpunkten in einzelnen Ländern.

Darüber hinaus kann im erweiterten Kreis eines multinationalen Konsortiums von „In-Depth“ Projekten eine Kommunikation entstehen, um auf bewährte Prinzipien und Erfahrungen aus dem Bereich der Verkehrssicherheit in anderen Ländern zurückgreifen zu können, mit dem Ziel, Stellhebel zur Verbesserung der globalen Verkehrssicherheit schnell und effektiv zu identifizieren. Nicht zu vernachlässigen ist auch, dass ein vereinfachter Zugang zu, und verbesserte Analysemöglichkeiten mit einem großen multinationalen Datenpool, treibende Kräfte für eine übergreifende Forschung auf dem Gebiet der Fahrzeugsicherheit sein können. Die jeweiligen „In-Depth“ Erhebungen vor Ort beinhal-

ten im Allgemeinen einen erheblich größeren Umfang an Variablen und auch innerhalb der Variablen eine größere Detailgenauigkeit. Für tiefer gehende Studien wird deshalb der jeweilige Originaldatensatz benötigt. Dadurch wird eine attraktive Plattform für die Organisationen der lokalen „In-Depth“ Erhebungen geboten, um Studien oder Daten anbieten zu können und damit die eigene Erhebung finanziell zu stärken.

IGLAD Standard Datensatz

Ausgangssituation für eine Harmonisierung sind verschiedene „In-Depth“ Datenerhebungen, die sich in Anzahl, Art und Ausprägungen der Faktoren voneinander unterscheiden. Für die Harmonisierung der Daten wurde ein gemeinsames Datenschema erzeugt, das als zusätzliche Schicht auf die unterschiedlichen Datenschemata aufgesetzt wird (Abbildung 2). Dieses gemeinsame Datenschema muss genau spezifiziert werden und bildet in Verbindung mit den in dieses Schema überführten Daten den Standarddatensatz.

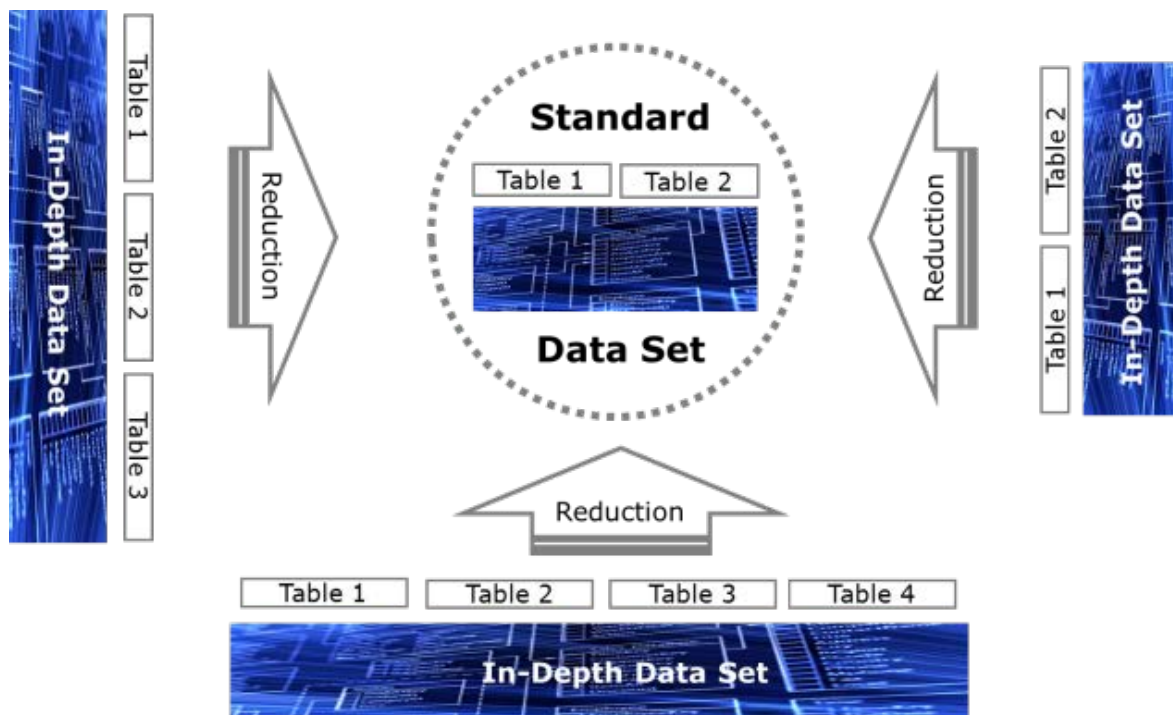


Abbildung 2: Standarddatensatz als zusätzliche Schicht auf dem „In-Depth“ Datensatz.

Die Anzahl und Ausprägungen der Faktoren des Standarddatensatzes sind im Mittelfeld zwischen nationaler Statistik und „In-Depth“ Datenerhebung anzusiedeln. Der Standarddatensatz beschränkt sich also auf die wesentlichsten Inhalte. Dies erleichtert die Erzeugung des Standarddatensatzes aus einer bereits vorhandenen „In-Depth“ Erhebung, bietet jedoch auf der anderen Seite nicht die Detailtiefe wie der ursprüngliche „In-Depth“ Datensatz. Sollte bei einer Analyse der Standarddatensatz nicht ausreichen,

muss immer noch auf den jeweiligen „In-Depth“ Datensatz zurückgegriffen werden, wobei der Standarddatensatz als Zwischenschicht die weitere Analyse erleichtern kann. Abbildung 3 gibt eine Übersicht über die Tabellen und Variablen des ersten IGLAD Codebooks. Insgesamt sind es 75 Variablen, die sich auf vier Tabellen verteilen.

Accident	Participant	Occupant	Safety System
Time Location Description Sketch Collision Type Accident Type Road Type Road Condition Light Condition Emergency	Vehicle Make Vehicle Model Registration Year Vehicle Mass Engine Type Number of Seats Primary/Secondary Collision with CDC, Initial Speed, Deceleration, Angle Collision Speed Delta V EES	Age Gender Weight Height Injury Severity MAIS AIS Regions	Type Use Deployment/ Activation

4 Tabellen, 75 Variablen

Abbildung 3: Datenschema der IGLAD Datenbank

Ohne die Verwendung eines Standarddatensatzes stellt eine vergleichende Auswertung verschiedener „In-Depth“ Datensätze einen erheblichen Aufwand dar, der zusätzlich sehr fehleranfällig ist, da oft die genauen Kontextinformationen der Kodierung der Daten dem Datenanalysten nicht geläufig sind.

Wird hingegen jeder einzelne Datensatz zunächst auf eine einheitliche Form gebracht, fällt dieser Aufwand nur einmal an und der so resultierende umfassende Datensatz kann leicht analysiert werden. Hinzu kommt, dass die Transformation des einzelnen Datensatzes in den standardisierten Datensatz von den Experten erzeugt wird, die die Daten auch erheben, was zu qualitativ besseren und konsistenteren Ergebnissen führt. Insbesondere können so regionale Unterschiede bei der Interpretation der Faktoren genauer und mit dem richtigen Kontextwissen berücksichtigt werden.

Falls eine nationale Statistik zur Verfügung steht, und ausreichend gemeinsame Faktoren mit dem „In-Depth“ Datensatz vorhanden sind (z.B. Ortslage, Unfallschwere, Art der Beteiligung, Tageszeit, Unfalltyp), können die Randverteilungen der gemeinsamen Faktoren mit geeigneten statistischen Methoden angeglichen werden, wie z.B. dem „Iterati-

ve Proportional Fitting“ [7]. Daraus können Wichtungsfaktoren berechnet werden, mit denen die gesamte Stichprobe an den Standarddatensatz angepasst werden kann. Stellt der „In-Depth“ Datensatz bereits eine eher repräsentativ gezogene Stichprobe dar (bezogen auf die nationale Statistik als Gesamtheit), kann dieser Schritt nahezu überflüssig werden. Eigenschaften einer Stichprobe, die die Repräsentativität bzgl. der nationalen Statistik begünstigen, sind z.B. identische Selektionskriterien (z.B. alle Unfälle mit Personenschaden), ein guter Querschnitt der geographischen und infrastrukturellen Gegebenheiten der Gesamtregion (des Landes) oder eine Zufallsstichprobe (z.B. durch zeitliche Stichprobenziehung mit alternierenden Schichten). Ferner sollte die Größe der Stichprobe ausreichend sein. Für IGLAD wurde eine Untergrenze von 50 Fällen und eine Obergrenze von 200 Fällen pro Jahr für jedes Land festgelegt. Durch ein festgelegtes Sample-Verfahren wird aus dem jeweiligen Originaldatensatz des Landes eine Stichprobe für IGLAD gezogen.

Der konsistenten Qualität der Daten wird hohe Priorität eingeräumt: Daten, die in IGLAD einfließen, müssen einen gewissen Qualitätsstandard erfüllen (z.B. maximale Unbekanntrate) und werden einer Reihe von Plausibilitätschecks unterzogen. Ferner werden bereits vorhandene internationale Kodierungsstandards verwendet (AIS, CDC).

Beim Aufbau neuer In-Depth Datenerhebungen können durch frühzeitige Einbindung in das IGLAD Konsortium Erfahrungswerte aus diesem mit einfließen. Dadurch gewinnen auch aus technischer Perspektive neue Erhebungen einen Vorsprung durch bereits vorhandene Dokumentation (Datenschema) und Expertise aus dem Konsortium.

Historie und Status

Initiiert von der Daimler AG, ACEA und verschiedenen Forschungsinstituten wurde das IGLAD Projekt zunächst als Arbeitsgruppe der FIA Mobility Group im Oktober 2010 ins Leben gerufen. Vorbereitende Arbeiten zur Erzeugung eines ersten Datensatzes einer internationalen In-Depth Unfalldatenbank fanden in 2012 statt, wobei eine Studie von FIA/CEESAR zur Ermittlung der weltweit verfügbaren Unfalldatenbanken durchgeführt wurde. Anschließend wurden in einer Pilotstudie von ACEA Daten aus acht Ländern (USA, Indien, Deutschland, Schweden, Frankreich, Spanien, Österreich und Polen) als „Proof of Concept“ in das Standarddatenschema überführt. In 2013 startete Phase 1 des Projektes mit dem Ziel, einen ersten vollständigen Datensatz zu erzeugen. Erste Schätzungen des Aufwandes machten schnell deutlich, dass eine Anschubfinanzierung notwendig war. Ein entsprechender Antrag wurde bei ACEA gestellt und genehmigt, mit der Auflage, dass das Projekt langfristig in der Lage ist sich selbst zu finanzieren. Um dies zu ermöglichen, wurde ein Finanzierungsmodell entwickelt, das IGLAD in die Lage versetzen sollte in Zukunft ohne Drittmittel auszukommen. Eine Arbeitsgruppe wurde gebildet, um eine ausgewogene Lösung zu finden, die die verschiedenen Rollen der

Partner und mögliche Kombinationen dieser in dem Projekt berücksichtigt. Auf das so entstandene Finanzierungsmodell wird näher in dem Abschnitt „Phase 2“ eingegangen. Weitere Einzelheiten über die Entwicklung des Projektes kann auch den Veröffentlichungen [8], [9], [10] entnommen werden.

IGLAD Projektphase 1

Die von ACEA finanzierte Phase 1 des Projektes wurde in 2014 abgeschlossen. Während dieser Phase wurden folgende Ziele erreicht:

- Das gemeinsame Datenschema wurde definiert und in einem Codebook dokumentiert
- Ein Consortium Agreement wurde entworfen und von den beteiligten Projektpartnern unterschrieben
- Ein Sample-Verfahren zur repräsentativen Stichprobenziehung in den einzelnen Ländern wurde definiert
- Das Rekodieren der Daten in die initiale Datenbank wurde abgeschlossen
- Die initiale Datenbank wurde am 6. Juni 2014 veröffentlicht



Abbildung 4: Phase 1 Datensatz mit 1550 Unfällen aus 10 Ländern (USA, Indien, Australien, Deutschland, Schweden, Frankreich, Spanien, Tschechien, Österreich, Italien)

Abbildung 4 gibt einen Überblick über die Länder und Datenlieferanten für den initialen Datensatz der Phase 1. Der erste Datensatz wurde am 6. Juni 2014 veröffentlicht mit folgenden Eckdaten:

1.550	Dokumentierte Unfälle
2.875	Beteiligte Fahrzeuge und Fußgänger/Fahrradfahrer
3.902	Insassen von Fahrzeugen und Fußgänger/Fahrradfahrer
10.736	In den Fahrzeugen verbaute Sicherheitssysteme
2.022	Unfallskizzen in 1.450 MB

Nach Veröffentlichung der ersten Daten wurden verschiedene erste Analysen durchgeführt. Ein Schwerpunkt war die Überprüfung der Repräsentativität des so erzeugten Datensamples bzgl. der nationalen Daten der jeweiligen beteiligten Länder. Als Vergleichsbasis dienten die Zahlen des OECD Projektes IRTAD [5], das vereinheitlichte Zahlen über Verkehrsunfälle in ausgewählten Ländern weltweit zusammenträgt und vereinheitlicht darstellt. Der Fokus von IRTAD liegt auf Getöteten im Straßenverkehr, was eine Untersuchung der Repräsentativität etwas erschwert, da der Anteil der Getöteten in IGLAD noch recht gering ist. Eine erste Analyse erbrachte jedoch recht interessante Ergebnisse. Die Abbildungen 5 und 6 zeigen einzelne Länder im Vergleich, aufgeteilt in Diagramme für bestimmte Teilpopulationen, wie Altersgruppen in Abbildung 5 und Verkehrsteilnahme in Abbildung 6.



Abbildung 5: Verteilung der Altersgruppen in einzelnen Ländern in IGLAD und in der nationalen Statistik.

Die blaue Linie zeigt den jeweiligen prozentualen Anteil der Teilpopulation in IGLAD und die rote Linie den in der offiziellen nationalen Statistik aus IRTAD. Liegen die beiden Linien übereinander, besteht eine sehr gute Übereinstimmung und IGLAD kann als re-

präsentative Stichprobe bzgl. der nationalen Statistik und dem dargestellten Merkmal angesehen werden.

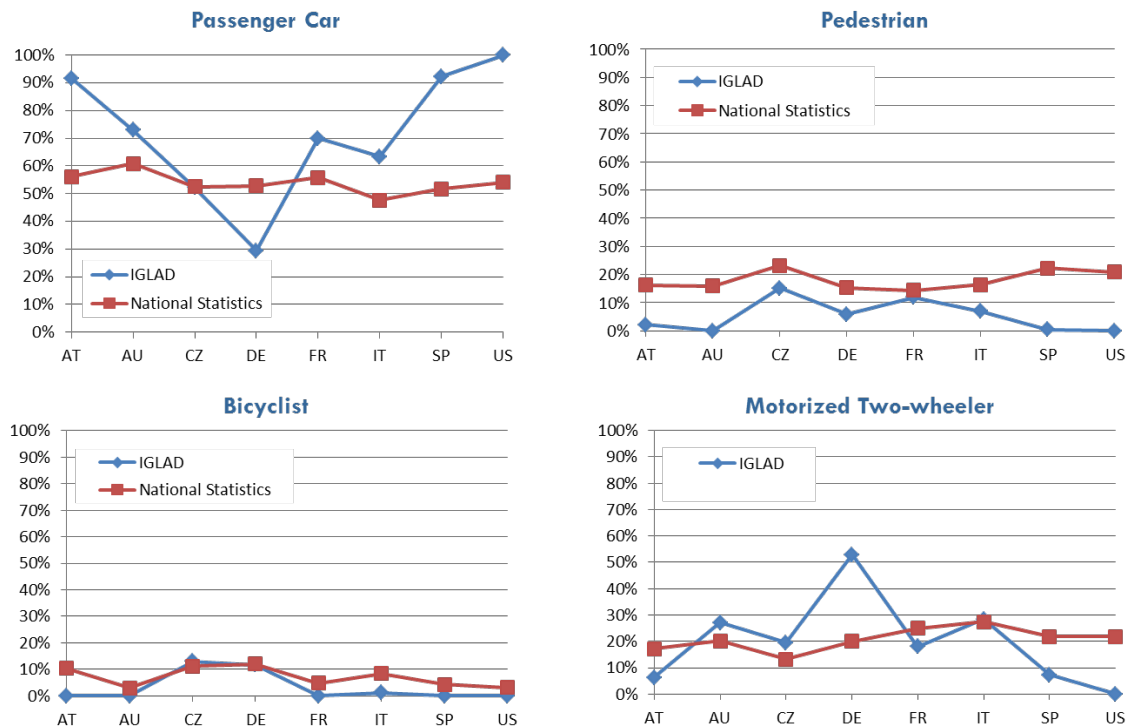


Abbildung 6: Verteilung der Verkehrsbeteiligung in einzelnen Ländern in IGLAD und in der nationalen Statistik.

Generell ist eine gute Übereinstimmung zwischen IGLAD und der nationalen Statistik zu beobachten. Größere Abweichungen gibt es bei der Verkehrsbeteiligung PKW. Dies kommt nicht unerwartet, da von einigen Ländern bereits bekannt war, dass die IGLAD Stichprobe Verschiebungen aufweist, da die ursprünglichen Daten bereits nach bestimmten Kriterien vorselektiert wurden. Für Datenanalysen mit entsprechendem Fokus steht dann immer noch die Möglichkeit der Wichtung offen, die durch in IGLAD erhobene Variablen aus der nationalen Statistik ermöglicht wird. Der Datensatz der Phase 1 soll zudem einen Startpunkt darstellen um sukzessive die einzelnen Datensätze der Länder zu verbessern, sowohl was Repräsentativität angeht, als auch im Hinblick auf Qualität und Quantität.

IGLAD Projektphase 2

Phase 2 des Projektes startete in 2014 im Anschluss an Phase 1. Da auf keine Fremdfinanzierung zurückgegriffen werden konnte, musste eine neue Projektstruktur und ein Finanzierungsmodell etabliert werden. Die Herausforderung bestand darin, eine große Anzahl von Organisationen innerhalb einer Projektstruktur zusammenzuführen, ohne dass eine Dachorganisation existiert, die auch die Finanzierung übernimmt. Dazu musste ein Modell entwickelt werden, das einen ausgeglichenen Fluss von Daten und Geldmitteln regelte. Ferner war eine Instanz notwendig, die verschiedene administrative

Aufgaben übernehmen konnte. Die Verantwortlichkeiten dieses Administrators sind im Einzelnen:

- Verwalten der finanziellen Ressourcen und Führen dedizierter Konten
- Vertragsabwicklung mit den Projektpartnern (Mitglieder und Datenlieferanten).
- Rechnungsstellung an Mitglieder des Projektes
- Auszahlungen an Datenlieferanten, Pufferung von Projektbudget, Anpassung von Zahlungen um ein ausgeglichenes Budget zu gewährleisten
- Erstellen regelmäßiger Finanzberichte
- Zusammenführen der Datensätze von den Datenlieferanten in einen auszuliefernden IGLAD Datensatz
- Verwalten eines Webspace um Daten und Projektdokumente vorzuhalten, Verwaltung von Mitglieder-Accounts

Zur Koordination der Projektpartner und Weiterentwicklung des Projektes war die Bildung eines Konsortiums (Steering Group) notwendig, in dem die Mitglieder des Projektes vertreten sind. Das IGLAD Konsortium sollte sowohl den Bestand des Datenpools sichern und organisatorische und technische Probleme adressieren, als auch eine Informationsplattform für beteiligte Institutionen darstellen, und so zu einem leichterem Austausch von Erfahrungen und Best Practices führen. Die einzelnen Aufgaben des Konsortiums sind:

- Generelle Angelegenheiten, Rechte und Pflichten: Konstitution und Erneuerung des Konsortiums. Klärung rechtlicher Fragen. Definition der Hauptziele und Rollen innerhalb des Projektes und Durchsetzung dieser. Erstellen von Statusberichten. Öffentlichkeitsarbeit und Vereinbarungen mit externen Organisationen.
- Regelung der Mitgliedschaften: Aufnahme und Ausschluss von Mitgliedern, Richtlinien für die Datennutzung, Kommunikation mit den Projektpartnern.
- Verwaltung der Datenbasis: Beauftragungen und Dienstleisterverträge, Regelung technischer Probleme, Qualitätskontrolle, Akquisition neuer Datenquellen, Definition und Regelung des Prozesses zur Datenverteilung.
- Gremien, Organisation und Meetings: Definition der Untergremien (Technical WG, Steering Group, Assembly), Organisation der Steering Group Meetings, Erstellen von Protokollen und Aufgabenlisten, Verteilung von Dokumenten.

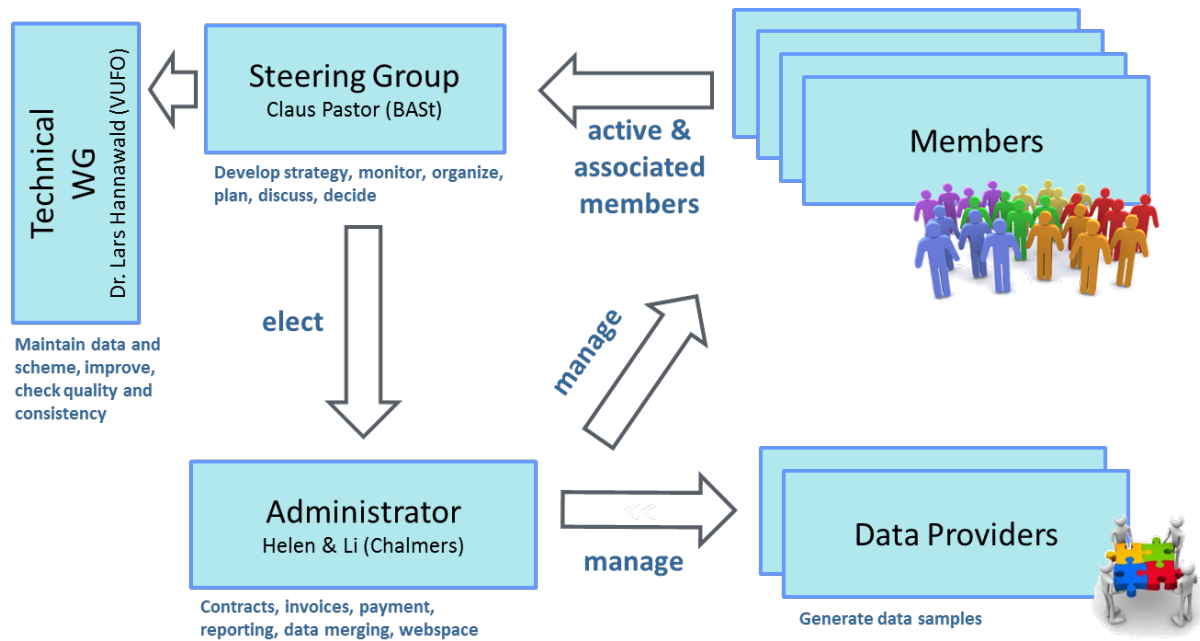


Abbildung 7: Phase 2 Projektstruktur

Zusätzlich zu dem Administrator und der Steering Group wurde eine technische Arbeitsgruppe (Technical WG) etabliert zur Klärung technischer Detailfragen. Diese Gruppe berichtet in die Steering Group und setzt sich überwiegend aus Vertretern der Datenlieferanten zusammen. Aufgaben dieser Gruppe sind im Einzelnen:

- Pflege und Weiterentwicklung des gemeinsamen Datenschemas und Codebooks
- Klärung technischer Anfragen und Bereitstellen von Hintergrundinformationen
- Ansprechpartner für Rekodierer
- Sammeln von Best Practices und Richtlinien
- Bereitstellen gemeinsamer Werkzeuge (Software-Tools)
- Entwicklung geeigneter Methoden zum Datenaustausch
- Anwenden und Erstellen von Integritäts-, Plausibilitäts- und Qualitätschecks
- Stichproben- und Hochrechnungstechniken

Abbildung 7 zeigt einen Überblick über die Organisationsstruktur mit den beschriebenen Gruppen und ihren Beziehungen untereinander. Als Administrator wurde die Chalmers University von der Steering Group gewählt, so dass die separate Gründung einer entsprechenden Organisationsform nicht notwendig war. Den Vorsitz über die Technical WG übernahm Dr. Hannawald, Leiter der VUFO GmbH in Dresden. Für die Leitung der Steering Group konnte Hr. Pastor von der BAST gewonnen werden.

Das Modell zur Finanzierung von IGLAD in der Projektphase 2 wurde lange diskutiert. Schließlich sollte es den Fortbestand des Projektes ermöglichen und es sollte gleichzeitig die verschiedenen Interessen der Projektpartner widerspiegeln und ausgleichen. Ein

erster Gedanke und zudem sehr einfacher Ansatz wäre, die Daten einfach zwischen den einzelnen Datenlieferanten auszutauschen, ohne finanzielle Mittel zu involvieren. Leider erwies sich dieser Ansatz recht schnell als wenig praktikabel, da viele Projektpartner entweder nur Daten für Analysen nutzen bzw. nur für die Verwendung im Projekt erzeugen wollten. Darüber hinaus gab es allein aufgrund der Größe der Organisationen Unterschiede zwischen den Projektpartnern, die es auszugleichen galt: während es für kleinere Datenlieferanten schwierig ist, die geforderte Mindestanzahl von Fällen pro Jahr zu erfüllen, haben größere Datenlieferanten keine Schwierigkeiten ein Sample aus ihren Daten zur Verfügung zu stellen. Hinzu kommt, dass größere Projektpartner, meist von mehreren Organisationen finanziert werden, welche dann auch von einem geteilten Datensatz profitieren würden. All dies führt zu unverhältnismäßig großen Unterschieden zwischen kleinen und großen Projekten wenn es um Nutzen und Beiträgen zu dem Projekt geht. Die Lösung dieses Problems liegt in der sauberen Trennung der Rollen und Interessen und entsprechendem Ausgleich auf finanzieller Basis. Eine zentrale Randbedingung ist, dass das Finanzierungsmodell ausschließlich das Fortbestehen des Projektes gewährleistet soll, um Daten für Forschungszwecke zur Verfügung zu stellen und dabei keinerlei Profit zu erzeugen. Diese Präambel ist auch fest im IGLAD Consortium Agreement verankert.

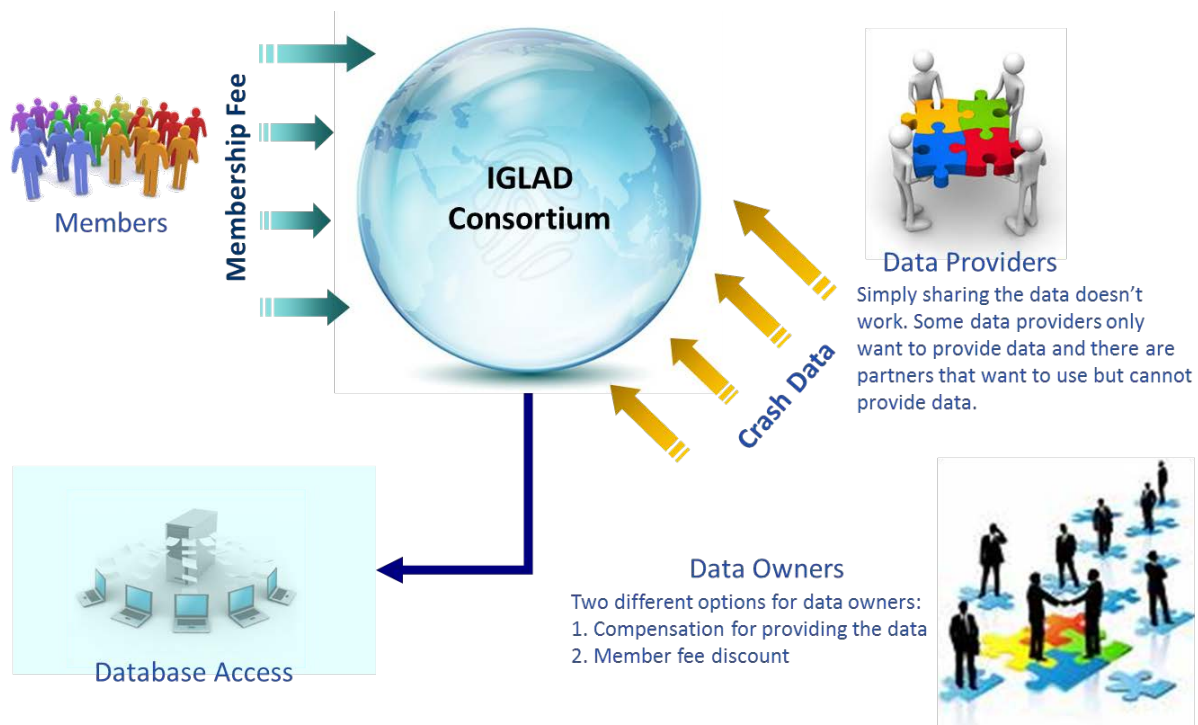


Abbildung 8: IGLAD Phase 2 Finanzierungsmodell

Entsprechend der Anforderungen wurden verschiedene Finanzierungsmodelle entwickelt. Dabei wurden drei verschiedenen Rollen im Projekt identifiziert und separiert: Mitglieder (Members), Datenlieferanten (Data Providers) und Dateneigner (Data Owners), was in Abbildung 8 dargestellt ist. Jedem Projektpartner kann eine oder auch mehrere dieser Rollen zugeordnet werden. Es ist sogar möglich, dass ein Projektpartner alle drei Rollen gleichzeitig innehat. Mitglieder sind Partner, die die Daten nutzen möchten, Datenlieferanten stellen dem Projekt Daten für ein bestimmtes Land zur Verfügung und Dateneigner müssen der Bereitstellung der Daten zustimmen.

In einem zweiten Schritt wurden Beziehungen zwischen den verschiedenen Rollen definiert und der Fluss von Daten und finanziellen Mitteln dargestellt. Um einen Ausgleich zwischen verschieden großen Projektpartnern zu schaffen, wurden zwei Optionen für die Dateneigner ermöglicht: Entweder ein fester finanzieller Ausgleich für das Bereitstellen der Daten oder bei großen Datenlieferanten eine Reduktion des Mitgliedsbeitrags für alle Organisationen, die diesen Datenlieferanten finanzieren (und damit Dateneigner sind).

Um sicherzustellen, dass das so definierte Finanzierungsmodell auch funktionieren würde, wurde eine Simulation durchgeführt, die auf Umfragen und „Letters of Intent“ basierte. Mit Start der Phase 2 konnte das Ergebnis dieser Simulation bestätigt werden, die Finanzierung für den ersten Datensatz der Phase 2 (Mitgliedsjahr 2014) war gesichert. Ende 2014 wurde bereits mit der Planung des zweiten Datensatzes der Phase 2 (Mitgliedsjahr 2015) begonnen. Eine kontinuierliche Fortsetzung der Simulation des Finanzierungsmodells unter Einberechnung der Veränderungen bei Mitgliedern und Datenlieferanten und einer möglichen Erweiterung des Datensatzes um Variablen aber auch Bildmaterial führt zu „worst“ und „best case“ Szenarien und ermöglicht eine langfristige Planung. Zum jetzigen Zeitpunkt ergibt sich eine positive Bilanz und die weitere Arbeit in dem Projekt scheint gesichert.

Mit Phase 2 wurden auch einige technische Verbesserungen umgesetzt. Um eine konsistente Qualität innerhalb des Datenpools zu ermöglichen, wurden einheitliche Plausibilitätschecks und weitere Prüfungsmechanismen eingeführt. Es wurde ein Bugtracker für Fehler- und Qualitätsmanagement installiert. Zur effizienteren Abwicklung der Datenerfassung und Vereinfachung der Zusammenführung der Daten bei gleichzeitiger Qualitätskontrolle, wurde ein einheitliches Software Tool zur Datenerfassung und für Qualitäts- und Konsistenzchecks bei den Data Providern installiert.

Die Rekodierung für das Mitgliedsjahr 2014 der Phase 2 wurde Ende 2014 gestartet und mit einer Datenlieferung ist Mitte 2015 zu rechnen. Mitglieder der Phase 2 sind: AIT, Autoliv, BAST, Bosch, BRSI, BMW, CASR, Chalmers, Daimler, Fiat, Ford, Nissan, Opel, PSA Renault, Takata, TATA, Toyota Europe, TU Graz, Uni Firenze, VCC, VUFO und VW.

Ausblick und nächste Schritte

Während bisher organisatorische Aspekte im Vordergrund standen, um die Grundlagen für einen ersten Datensatz zu legen, eine geeignete Infrastruktur für das Projekt zu schaffen und eine nachhaltige Finanzierung zu ermöglichen, verschiebt sich der Fokus mehr hin zu inhaltlicher Arbeit mit den Daten und gleichzeitig einer Erweiterung und Verbesserung der Daten.

So konnte Anfang 2015 ein flankierendes Projekt abgeschlossen werden, das genauere Informationen über die Pre-Crash Phase der Unfälle in den IGLAD Daten bereitstellen sollte. Diese „Pre-Crash Matrix“ umfasst eine Trajektorie von jedem Beteiligten bis zu 5s vor Kollision und zugehörige fahrdynamische Werte. Dafür wurde zunächst untersucht, welche Informationen notwendig sind, um eine Simulation des Unfalles durchführen zu können, die dann die gewünschten Daten erzeugen würde. In einem nächsten Schritt wurde eine Bestandsaufnahme gemacht, um festzustellen, welche Datensätze aus IGLAD geeignet für eine solche Simulation sind. In einer Pilotstudie wurden dann für jedes in IGLAD vertretene Land fünf Unfälle simuliert und eine Pre-Crash Matrix erstellt.

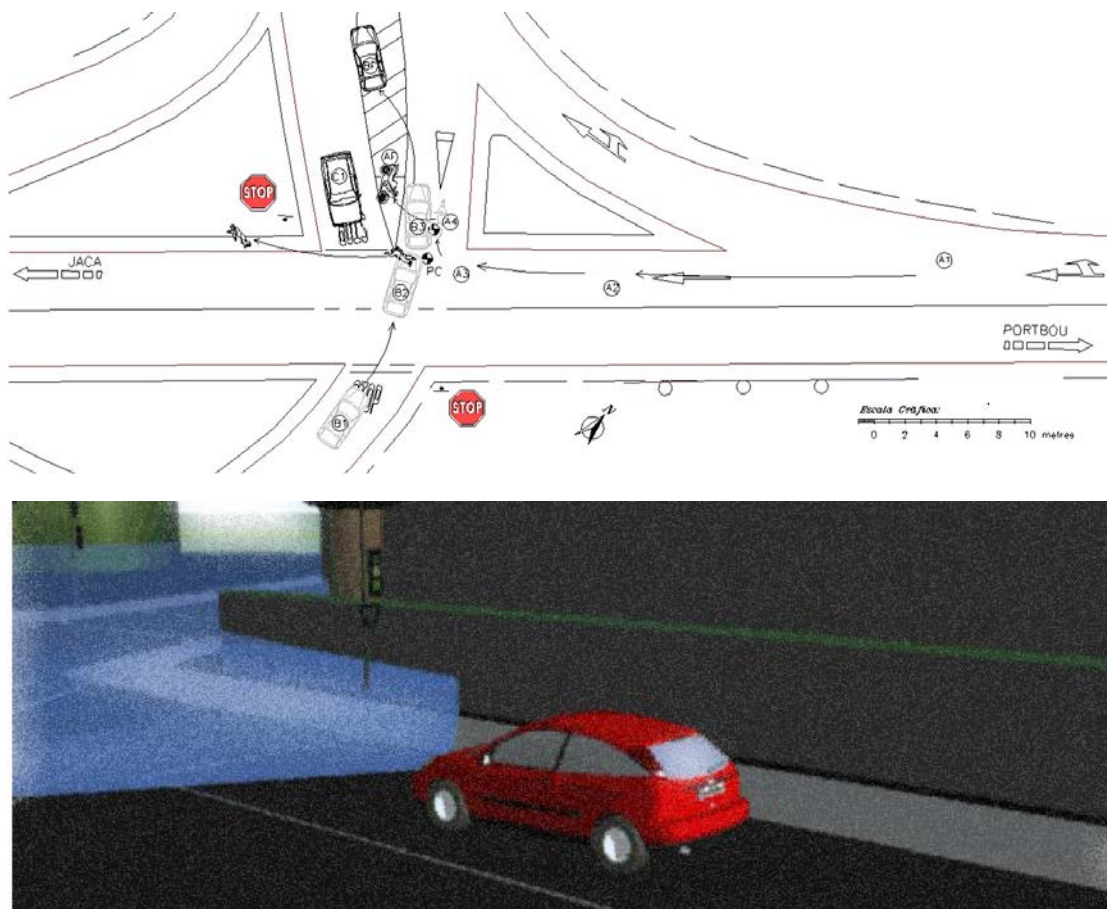


Abbildung 9: Simulation der Pre-Crash Phase

Die so erzeugten Daten über die Pre-Crash Phase des Unfalls ermöglichen die Simulation von Sicherheitssystemen und eine Bewertung von deren Performance in realen Unfällen in allen bei IGLAD beteiligten Ländern. Abbildung 9 zeigt die Skizze eines Unfalles aus IGLAD, die die wesentliche Grundlage für die Erstellung der Pre-Crash Matrix darstellt, und die entsprechende Simulation mit einem im Fahrzeug verbauten Assistenzsystem für Kreuzungssituationen. Ein nützlicher Nebeneffekt dieser Arbeiten ist die Verbesserung der Rekonstruktionsdaten und Skizzen der Unfälle aus IGLAD.

Für die kommenden Veröffentlichungen von Datensätzen sind ebenfalls Erweiterungen geplant, sowohl bei der Anzahl der beteiligten Länder, die Daten zur Verfügung stellen, als auch beim Umfang der Daten. So wird z.B. über eine Bereitstellung von Fotodokumentation in Ergänzung zu den schon vorhandenen Unfallskizzen nachgedacht. Weitere aktuelle Informationen über das Projekt stehen über die Webseite [11] zur Verfügung.

Zusammenfassung

Die globale Verkehrssicherheit rückt mit dem stark wachsenden Automobilmarkt in den Schwellenländern immer mehr in den Vordergrund. Während in technologisch und infrastrukturell weiter entwickelten Ländern die Anzahl der Getöteten im Straßenverkehr kontinuierlich rückläufig sind, ist in Schwellenländern ein entgegengesetzter Trend zu beobachten. Das Projekt IGLAD adressiert diese Problematik aus Sicht der Unfallforschung, in dem die Datenbasis von Straßenverkehrsunfällen weltweit harmonisiert und vergrößert werden soll. Dies soll dazu führen, dass die Stellhebel zur Verbesserung der globalen Verkehrssicherheit schnell und effektiv identifiziert werden können. Das Projekt konnte erfolgreich gestartet werden, ein erster Datensatz liegt vor und der zweite befindet sich in Vorbereitung. Auch organisatorisch und finanziell konnte eine Grundlage aufgebaut werden, die ein nachhaltiges Fortbestehen des Projektes sichern soll. Die nächsten Schritte werden eine erweiterte Nutzung und Verbesserung der Datenbasis sein. Insgesamt konnte das Projekt die gesteckten Ziele zügig umsetzen und es bleibt zu hoffen, dass der weitere Fortschritt in gleichem Maße erfolgen kann und damit eine ausreichend große Datenbasis in naher Zukunft aufgebaut werden kann, um Fragen der globalen Verkehrssicherheit besser adressieren zu können.

Literatur

- [1] M. Peden et. al., WHO, World report on road traffic injury prevention, 2004.
- [2] E. Kopits, M. Cropper, Traffic fatalities and economic growth. Washington, DC, The World Bank, 2003 (Policy Research Working Paper No. 3035).
- [3] EU Commission, European Road Safety Action Programme, ISBN 92-894-5893-3, 2003.
- [4] Euro NCAP Advanced reward, <http://www.euroncap.com/rewards.aspx>, 2010.
- [5] IRTAD, International Traffic Safety Data and Analysis Group, International Transport Forum, OECD, <http://internationaltransportforum.org/Irtadpublic/index.html>.

- [6] CARE, European Road Safety Observatory, Safetynet (EU Commission DG-TREN), http://erso.swov.nl/safetynet/content/wp_1_care_accident_data.htm.
- [7] W.E. Deming, F.F. Stephan, On a Least Squares Adjustment of a Sampled Frequency Table When the Expected Marginal Totals are Known, *Annals of Mathematical Statistics* 11 (4): 427–444. doi:10.1214/aoms/1177731829. MR3527, 1940.
- [8] Ockel, Bakker, Schöneburg, “Internationale Harmonisierung von Unfalldaten”, VDI Kongress 2011, Berlin
- [9] Bakker, “iGLAD - A pragmatic approach for an international accident database”, VDA Kongress 2012, Sindelfingen
- [10] Ockel, Bakker, Schöneburg, “An initiative towards a simplified international in-depth accident database”, ESAR Conference 2012, Hannover
- [11] <http://www.iglad.net>