

Karin Edvardsson Björnberg
Sven Ove Hansson
Matts-Åke Belin
Claes Tingvall
Editors

The Vision Zero Handbook

Theory, Technology and Management
for a Zero Casualty Policy

OPEN ACCESS

 Springer

The Vision Zero Handbook

Karin Edvardsson Björnberg • Sven Ove
Hansson • Matts-Åke Belin • Claes Tingvall
Editors

The Vision Zero Handbook

Theory, Technology and Management
for a Zero Casualty Policy

With 219 Figures and 62 Tables

 Springer

Editors

Karin Edvardsson Björnberg
Division of Philosophy
KTH Royal Institute of Technology
Stockholm, Sweden

Sven Ove Hansson
Division of Philosophy
Royal Institute of Technology (KTH)
Stockholm, Sweden

Department of
Learning, Informatics, Management
and Ethics
Karolinska Institutet
Stockholm, Sweden

Matts-Åke Belin
Department of Social
Determinants of Health WHO
Geneva, Switzerland

Claes Tingvall
Chalmers University of Technology
Gothenburg, Sweden

Monash University Accident Research Centre
Melbourne, Victoria, Australia

AFRY
Stockholm, Sweden



ISBN 978-3-030-76504-0

ISBN 978-3-030-76505-7 (eBook)

<https://doi.org/10.1007/978-3-030-76505-7>

© The Editor(s) (if applicable) and The Author(s) 2023. This book is an open access publication.

Open Access This book is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this book are included in the book's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the book's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

The use of general descriptive names, registered names, trademarks, service marks, etc. in this publication does not imply, even in the absence of a specific statement, that such names are exempt from the relevant protective laws and regulations and therefore free for general use.

The publisher, the authors and the editors are safe to assume that the advice and information in this book are believed to be true and accurate at the date of publication. Neither the publisher nor the authors or the editors give a warranty, expressed or implied, with respect to the material contained herein or for any errors or omissions that may have been made. The publisher remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.

This Springer imprint is published by the registered company Springer Nature Switzerland AG.
The registered company address is: Gewerbestrasse 11, 6330 Cham, Switzerland.

Preface

Globally, about 1.3 million people die every year in road traffic crashes and about 50 million are injured. For long, the death toll on roads was considered to be a necessary price that we have to pay for our mobility and development. But beginning in the late 1990s, an alternative approach to road safety has become more and more influential. The Vision Zero movement declares that every severe crash in road traffic is an avoidable failure and that no other goal is satisfactory than zero fatalities and serious injuries. This is by no means a new idea. Similar views have been expressed in many other areas of safety management. Safety work based on the idea that every accident is one too much can be found in workplace safety, fire safety, aviation, suicide prevention, patient safety, infection control, and many other areas. Zero goals are also gaining traction in environmental protection and sustainable development work. Terms such as zero waste, zero emissions, zero carbon, and zero poverty have become increasingly important in climate and environmental policies.

Detractors claim that Vision Zero is too stringent and therefore also unrealistic. But practical experience has shown again and again that the Vision Zero approach can efficiently reduce the number of deaths and serious injuries. Of course, it is not enough to set up Vision Zero as a goal. Its effects materialize when it is systematically applied, and every serious accident is treated as a failure that must not be repeated.

This handbook is the first comprehensive collection of knowledge and experience of Vision Zero. Its contributing authors are experts from all around the world, representing a wide range of academic disciplines and an equally wide range of specializations in practical safety management. The handbook is divided into five main parts. The *first part* discusses Vision Zero from several theoretical perspectives, relating it to other road safety targets, to other principles of safety management, to other forms of policy-making, and to ethical issues such as responsibility and paternalism. It also contains a chapter that scrutinizes the most common arguments against Vision Zero. The *second part* provides a broad overview of the worldwide adoption of Vision Zero in road traffic. It contains chapters on Vision Zero in Sweden, Norway, the Netherlands, Germany, Poland, Lithuania, Australia, Canada, the United States, and India, as well as chapters on its impact in international cooperation. The *third part* discusses Vision Zero from a managerial point of view, providing perspectives from road managers, vehicle manufacturers, and

consumers. The *fourth part* introduces tools and technologies for Vision Zero in road traffic. It includes chapters on road safety analysis, road design, speed control, driver distraction, and automated vehicles. Finally, the *fifth part* puts focus on the application of Vision Zero in other areas than road traffic. It contains chapters devoted to Vision Zero in workplace safety, fire safety, suicide prevention, disease eradication, and waste management.

We hope that this handbook will inspire further developments, innovations, and decisions that contribute to eliminating human suffering. Safety is achievable, but it requires commitment, knowledge, and careful planning.

Stockholm, Sweden
Stockholm, Sweden
Geneva, Switzerland
Gothenburg, Sweden
November 2022

Karin Edvardsson Björnberg
Sven Ove Hansson
Matts-Åke Belin
Claes Tingvall

Contents

Part I Ideas and Principles	1
1 Vision Zero and Other Road Safety Targets	3
Karin Edvardsson Björnberg	
2 Zero Visions and Other Safety Principles	31
Sven Ove Hansson	
3 Arguments Against Vision Zero: A Literature Review	107
Henok Girma Abebe, Sven Ove Hansson, and Karin Edvardsson Björnberg	
4 What Is a Vision Zero Policy? Lessons from a Multi-sectoral Perspective	151
Ann-Catrin Kristianssen and Ragnar Andersson	
5 Responsibility in Road Traffic	177
Sven Ove Hansson	
6 Liberty, Paternalism, and Road Safety	205
Sven Ove Hansson	
Part II Vision Zero: An International Movement for Traffic Safety	243
7 Vision Zero: How It All Started	245
Claes Tingvall	
8 Vision Zero in Sweden: Streaming Through Problems, Politics, and Policies	267
Matts-Åke Belin	
9 Vision Zero in Norway	295
Rune Elvik	
10 Sustainable Safety: A Short History of a Safe System Approach in the Netherlands	307
Fred Wegman, Letty Aarts, and Peter van der Knaap	

11	Vision Zero in Germany	337
	Wolfram Hell, Kurt Bodewig, Ute Hammer, Christian Kellner, Clemens Klinke, Matthias Mück, Martin Schreiner, Felix Walz, and Guido Zielke	
12	Vision Zero in Poland	359
	Kazimierz Jamroz, Aleksandra Romanowska, Lech Michalski, and Joanna Żukowska	
13	Vision Zero in Lithuania	399
	Vidas Žuraulis and Vidmantas Pumputis	
14	Vision Zero in EU Policy: An NGO Perspective	439
	Ellen Townsend and Antonio Avenoso	
15	The Development of the “Vision Zero” Approach in Victoria, Australia	475
	Samantha Cockfield, David Healy, Anne Harris, Allison McIntyre, and Antonietta Cavallo	
16	Vision Zero on Federal Level in Canada	507
	Pamela Fuselli	
17	Adoption of Safe Systems in the United States	553
	Jeffrey P. Michael, Leah Shahum, and Jeffrey F. Paniati	
18	Establishing Vision Zero in New York City: The Story of a Pioneer	571
	Ann-Catrin Kristianssen	
19	Traffic Safety in India and Vision Zero	597
	Geetam Tiwari and Dinesh Mohan	
20	Vision Zero in the United Nations	637
	Meleckidzedek Khayesi	
21	Towards a Potential Paradigm Shift: The Role of Vision Zero in Global Road Safety Policy Making	647
	Ann-Catrin Kristianssen	
Part III	Management and Leadership for Vision Zero	673
22	ISO 39001 Road Traffic Safety Management System, Performance Recording, and Reporting	675
	Anders Lie and Claes Tingvall	
23	What the Car Industry Can Do: Volvo Cars	687
	Anders Eugensson and Jan Ivarsson	
24	What the Car Industry Can Do: Mercedes-Benz’ View	727
	Rodolfo Schöneburg and Karl-Heinz Baumann	

25	Consumer Ratings and Their Role in Improving Vehicle Safety	755
	Michiel R. van Ratingen	
26	Saving Lives Beyond 2020: The Next Steps	789
	Claes Tingvall, Jeffrey P. Michael, Peter Larsson, Anders Lie, Maria Segui-Gomez, Shaw Voon Wong, Olive Kobusingye, Maria Krafft, Fred Wegman, Margie Peden, Adnan Hyder, Meleckidzedeck Khayesi, Eric Dumbaugh, Samantha Cockfield, and Alejandro Furas	
27	Miscommunications Based on Different Meanings of “Safe” and Their Implications for the Meaning of Safe System	841
	Chika Sakashita, R. F. Soames Job, and Matts-Åke Belin	
Part IV	Tools and Technologies for Vision Zero	855
28	Road Safety Analysis	857
	Matteo Rizzi and Johan Strandroth	
29	Speed-Limits in Local Streets: Lessons from a 30 km/h Trial in Victoria, Australia	881
	Brian N. Fildes, Brendan Lawrence, Luke Thompson, and Jennie Oxley	
30	Urban Road Design and Keeping Down Speed	903
	Bruce Corben	
31	Rural Road Design According to the Safe System Approach	947
	Helena Stigson, Anders Kullgren, and Lars-Erik Andersson	
32	Speed and Technology: Different Modus of Operandi	971
	Matts-Åke Belin and Anna Vadeby	
33	Driver Distraction: Mechanisms, Evidence, Prevention, and Mitigation	995
	Michael A. Regan and Oscar Oviedo-Trespalcacios	
34	Automated Vehicles: How Do They Relate to Vision Zero	1057
	Anders Lie, Claes Tingvall, Maria Håkansson, and Ola Boström	
Part V	Vision Zero in Other Areas	1073
35	Vision Zero in Workplaces	1075
	Gerard I. J. M. Zwetsloot and Pete Kines	
36	Suicide in the Transport System	1103
	Anna-Lena Andersson and Kenneth Svensson	

37	Vision Zero in Suicide Prevention and Suicide Preventive Methods	1117
	Danuta Wasserman, I. Tadić, and C. Bec	
38	Vision Zero on Fire Safety	1143
	Ragnar Andersson and Thomas Gell	
39	Vision Zero in Disease Eradication	1165
	Mark Rosenberg, Emaline Laney, and Claes Tingvall	
40	Zero-Waste: A New Sustainability Paradigm for Addressing the Global Waste Problem	1195
	Atiq Zaman	
Index		1219

About the Editors



Karin Edvardsson Björnberg

Division of Philosophy
KTH Royal Institute of Technology
Stockholm, Sweden

Karin Edvardsson Björnberg is associate professor of environmental philosophy in the Department of Philosophy and History at the Royal Institute of Technology (KTH), Stockholm. She has led several research projects on goal-setting in climate adaptation, delay mechanisms in environmental policy, and ethics of biodiversity off-setting. She takes an interest in risk and safety issues and has published articles on EU regulation of genetically modified crops and goal-setting in road traffic safety policy. She currently serves as scientific expert at the Swedish Ethical Review Authority and is vice director of third cycle education at the School of Architecture and the Built Environment (KTH).



Sven Ove Hansson

Division of Philosophy
Royal Institute of Technology (KTH)
Stockholm, Sweden

Department of
Learning, Informatics, Management
and Ethics
Karolinska Institutet
Stockholm, Sweden

Sven Ove Hansson is professor emeritus in philosophy at the Royal Institute of Technology (KTH), Stockholm.

He is a leading researcher in the philosophy of risk and safety. He is also active in other areas of philosophy, in particular the philosophy of science and technology, decision theory, logic, and moral and political philosophy. He is editor-in-chief of *Theoria* and of the two book series *Outstanding Contributions to Logic* and *Philosophy, Technology, and Society*. He is a member of the Royal Swedish Academy of Engineering Sciences and past president of the Society for Philosophy and Technology. Hansson has published 17 books and about 400 papers in refereed international journals and books. His books include *The Ethics of Risk* (2013), *The Ethics of Technology: Methods and Approaches* (edited, 2017), and *Responsibility for Health* (2022).



Matts-Åke Belin

Department of Social
Determinants of Health WHO
Geneva, Switzerland

Matts-Åke Belin is adjunct professor at the Royal Institute of Technology (KTH), Stockholm, and he has worked as senior policy advisor at the Swedish Transport Administration where he was responsible for the development of Vision Zero Academy. Since January 2022, Dr. Belin has joined the World Health Organization where he will be leading efforts by WHO to support member states and other global actors in the implementation of the Global Plan for the Decade of Action for Road Safety 2021–2030.



Claes Tingvall

Chalmers University of Technology
Gothenburg, Sweden

Monash University Accident Research Centre
Melbourne, Victoria, Australia

AFRY, Stockholm, Sweden

Claes Tingvall is adjunct professor at the Division of Vehicle Engineering and Autonomous Systems, Chalmers University of Technology, and at the Accident Research Centre, Monash University. He is also a Senior

Consultant at the engineering and design company AFRY. Previously, as Director of Traffic Safety at the Swedish Transport Administration, he had a decisive role in the creation of Vision Zero. He continues to have a leading role in international discussions on Vision Zero and traffic safety in general.

Contributors

Letty Aarts SWOV Institute for Road Safety Research, The Hague, The Netherlands

Henok Girma Abebe Division of Philosophy, KTH Royal Institute of Technology, Stockholm, Sweden

Anna-Lena Andersson Swedish Road Administration, Trollhättan, Sweden
Department of Orthopedics, Institute of Clinical Sciences at Sahlgrenska Academy, University of Gothenburg, Gothenburg, Sweden

Lars-Erik Andersson AFRY, Stockholm, Sweden

Ragnar Andersson The Centre for Societal Risk Research, CSR, Karlstad University, Karlstad, Sweden

Antonio Avenoso ETSC, Brussels Region, Belgium

Karl-Heinz Baumann SafetyFirst, Independent Scholar, Starzach Boerstingen, Germany

C. Bec National Centre for Suicide Research and Prevention of Mental Ill-Health (NASP), Karolinska Institutet, Stockholm, Sweden

Matts-Åke Belin Department of Social Determinants of Health WHO, Geneva, Switzerland

Kurt Bodewig Deutsche Verkehrswacht, Berlin, Germany

Ola Boström Veoneer AB, Stockholm, Sweden

Antonietta Cavallo Cavallo Road Safety Consulting, Melbourne, Victoria, Australia

Samantha Cockfield Transport Accident Commission (TAC), Geelong, Victoria, Australia

Bruce Corben Corben Consulting, Melbourne, VIC, Australia

Eric Dumbaugh Florida Atlantic University, Boca Raton, FL, USA

Karin Edvardsson Björnberg Division of Philosophy, KTH Royal Institute of Technology, Stockholm, Sweden

Rune Elvik Department of Safety and Security, Institute of Transport Economics, Oslo, Norway

Anders Eugensson Volvo Cars, Gothenburg, Sweden

Brian N. Fildes Monash University Accident Research Centre, Melbourne, VIC, Australia

Alejandro Furas Latin NCAP, Montevideo, Uruguay

Pamela Fuselli Parachute, Toronto, ON, Canada

Thomas Gell Centre for Societal Risk Research, Karlstad University, Karlstad, Sweden

Maria Håkansson AFRY (ÅF Pöyry AB), Stockholm, Sweden

Ute Hammer Deutscher Verkehrssicherheitsrat DVR, Bonn, Germany

Sven Ove Hansson Division of Philosophy, Royal Institute of Technology (KTH), Stockholm, Sweden

Department of Learning, Informatics, Management and Ethics, Karolinska Institutet, Stockholm, Sweden

Anne Harris Transport Accident Commission (TAC), Geelong, Victoria, Australia

David Healy Road Safety Consulting, Melbourne, Victoria, Australia

Wolfram Hell Gesellschaft für Medizinische und Technische Biomechanik e.V., gmtb (Society for Medical and Technical Biomechanics), Munich, Germany

Department of Forensic Epidemiology, Ludwig-Maximilians-University of Muenchen; gmtb, Muenchen, Germany

Adnan Hyder George Washington University, Washington, DC, USA

Jan Ivarsson Senior Technical Advisor Safety, Volvo Cars, Gothenburg, Sweden

Kazimierz Jamroz Gdansk University of Technology, Faculty of Civil and Environmental Engineering, Gdansk, Poland

R. F. Soames Job Global Road Safety Solutions Pty Ltd., Sydney, NSW, Australia

Christian Kellner Deutscher Verkehrssicherheitsrat DVR, Bonn, Germany

Meleckidzedek Khayesi Department of Social Determinants of Health, WHO, Geneva, Switzerland

Pete Kines National Research Centre for the Working Environment, Copenhagen, Denmark

Clemens Klinke Dekra Se, Stuttgart, Germany

Olive Kobusingye Makerere University, Kampala, Uganda

Maria Krafft Swedish Transport Administration, Borlänge, Sweden

Ann-Catrin Kristianssen School of Humanities, Education and Social Sciences, Örebro University, Örebro, Sweden

Anders Kullgren Folksam Insurance Group and Chalmers University of Technology, Stockholm, Sweden

Emaline Laney Brigham and Women's Hospital, Boston, MA, USA

Peter Larsson Swedish Transport Administration, Borlänge, Sweden

Brendan Lawrence Monash University Accident Research Centre, Melbourne, VIC, Australia

Anders Lie Swedish Ministry of Infrastructure, Chalmers University of Technology, Gothenburg, Sweden

Allison McIntyre Allison McIntyre Consulting, Melbourne, Victoria, Australia

Jeffrey P. Michael Center for Injury Research and Policy, Johns Hopkins University, Baltimore, MD, USA

Lech Michalski Gdansk University of Technology, Faculty of Civil and Environmental Engineering, Gdansk, Poland

Dinesh Mohan Indian Institute of Technology/Independent Council for Road Safety International, Delhi, India

Matthias Mück City of Munich, Munich, Germany

Oscar Oviedo-Trespalacios Centre for Accident Research and Road Safety-Queensland (CARRS-Q), Queensland University of Technology (QUT), Brisbane, QLD, Australia

Jennie Oxley Monash University Accident Research Centre, Melbourne, VIC, Australia

Jeffrey F. Paniati Institute of Transportation Engineers, Washington, DC, USA

Margie Peden The George University for Global Health, London, UK

Vidmantas Pumputis Department of Automobile Engineering, Faculty of Transport Engineering, Vilnius Gediminas Technical University, Vilnius, Lithuania

Michael A. Regan Research Centre for Integrated Transport Innovation (rCITI), School of Civil and Environmental Engineering, University of New South Wales, Sydney, NSW, Australia

Matteo Rizzi Swedish Transport Administration (STA), Borlänge, Sweden

Aleksandra Romanowska Gdansk University of Technology, Faculty of Civil and Environmental Engineering, Gdansk, Poland

Mark Rosenberg The Task Force for Global Health (President Emeritus), Atlanta, GA, USA

Chika Sakashita Global Road Safety Solutions Pty Ltd., Sydney, NSW, Australia

Rodolfo Schöneburg Road Safety Counselor (until 11/2021 Director for Vehicle Safety, Durability and Corrosion Protection at Mercedes-Benz), RSC Safety Engineering, Hechingen, Germany

Martin Schreiner City of Munich, Munich, Germany

Maria Segui-Gomez Bloomberg School of Public Health, Johns Hopkins University, Baltimore, MD, USA

School of Medicine, University of Virginia, Charlottesville, VA, USA

Leah Shahum Vision Zero Network, San Francisco, CA, USA

Helena Stigson Folksam Insurance Group, Chalmers University of Technology and Karolinska Institutet, Stockholm, Sweden

Johan Strandroth Swedish Transport Administration (STA), Borlänge, Sweden

Kenneth Svensson Swedish Road Administration, Gothenburg, Sweden

I. Tadić National Centre for Suicide Research and Prevention of Mental Ill-Health (NASP), Karolinska Institutet, Stockholm, Sweden

Luke Thompson Monash University Accident Research Centre, Melbourne, VIC, Australia

Claes Tingvall Chalmers University of Technology, Gothenburg, Sweden

Monash University Accident Research Centre, Melbourne, Victoria, Australia

AFRY, Stockholm, Sweden

Geetam Tiwari Department of Civil Engineering, Indian Institute of Technology, Delhi, India

Ellen Townsend ETSC, Brussels Region, Belgium

Anna Vadeby Swedish National Road and Transport Administration, VTI, Linköping, Sweden

Department of Mechanics and Maritime Sciences, Division of Vehicle Safety, Chalmers University of Technology, Göteborg, Sweden

Michiel R. van Ratingen European New Car Assessment Programme (Euro NCAP), Leuven, Belgium

Peter van der Knaap The Hague, The Netherlands

Felix Walz Gesellschaft für Medizinische und Technische Biomechanik e.V., gmttb (Society for Medical and Technical Biomechanics), Munich, Germany

Danuta Wasserman National Centre for Suicide Research and Prevention of Mental Ill-Health (NASP), Karolinska Institutet, Stockholm, Sweden

Fred Wegman Department of Transport and Planning, Delft University of Technology, Civil Engineering and Geosciences, Delft, The Netherlands

Shaw Voon Wong University Putra, Seri Kembangan, Malaysia

Atiq Zaman Curtin University Sustainability Policy Institute, School of Design and the Built Environment, Curtin University, Perth, WA, Australia

Guido Zielke Federal Ministry for Digital and Transport, Berlin, Germany

Joanna Żukowska Gdansk University of Technology, Faculty of Civil and Environmental Engineering, Gdansk, Poland

Vidas Žuraulis Department of Automobile Engineering, Faculty of Transport Engineering, Vilnius Gediminas Technical University, Vilnius, Lithuania

Gerard I. J. M. Zwetsloot Gerard Zwetsloot Research & Consultancy, Amsterdam, The Netherlands

Part I

Ideas and Principles

Vision Zero and Other Road Safety Targets

1

Karin Edvardsson Björnberg

Contents

Introduction	4
A Brief History of Goal-Setting in Road Safety Management	6
Why Use Quantified Goals in Road Safety Management?	8
When Is a Road Safety Target Rational?	10
Precision	11
Evaluability	13
Approachability	14
Motivity	15
Balancing the Criteria	17
Rationality Criteria for Systems of Goals	17
Completeness	18
The Number of Targets	19
Consistency	19
Why Does a Road Safety Target Have to Be Stable Over Time?	20
Who Should Be Involved in the Goal-Setting Process?	21
What About Contextual Rationality Aspects?	23
Conclusions	24
Cross-References	24
References	25

Abstract

Every year, around 1.3 million people are killed on the road and another 20–50 million are severely injured. This makes road safety one of the most critical global public health issues. To address the negative trend, the international community has responded with the adoption of road safety targets. Sustainable Development Goals 3.6 and 11.2 are two examples. Also at the national level, goals and targets are increasingly used to steer work towards improved road safety. The frequent use of goals and targets in road safety policy makes it interesting to investigate

K. Edvardsson Björnberg (✉)

Division of Philosophy, KTH Royal Institute of Technology, Stockholm, Sweden

e-mail: karin.bjornberg@abe.kth.se

under what conditions the adopted goals can be expected to be achieved. This chapter summarizes the main themes and conclusions of research that have been conducted on goal-setting in road safety policy and management to date. Drawing on previous research, it outlines and discusses a set of criteria that road safety goals should meet in order to be achievement-inducing, that is, have the capacity to guide and induce efforts towards the vision of zero fatalities and serious injuries on the road.

Keywords

Vision Zero · Goal-Setting · Management by Objectives · Rationality · Precision · Evaluability · Approachability · Motivity

Introduction

Every year, around 1.3 million people are killed on the road, and another 20–50 million are severely injured (WHO 2018). Data from the World Health Organization (WHO) show that around 90% of the fatalities occur in middle- and low-income countries. Approximately half of those killed are vulnerable road users, including pedestrians, cyclists and motorcyclists, children, and elderly people. At present, road traffic accidents are the dominant cause of death among people aged 5–29 years, and they are expected to be the seventh leading cause of death in 2030. Thus, road safety is one of the most critical global public health issues to date. In order to stabilize and then reverse the negative trend, the international community has responded with the adoption of road safety targets aimed at reducing the number of killed and seriously injured people on the road. In March 2010, the United Nations (UN) General Assembly proclaimed the period 2011–2020 as the “Decade of Action for Road Safety” (United Nations 2010). Since 2015, the UN member states have been bound by the Sustainable Development Goals (SDGs), which include the goal to halve the number of global deaths and injuries from road traffic accidents by 2020 (SDG 3.6; see also SDG 11.2). In November 2017, the UN agreed on a subset of global road safety targets designed to advance current efforts toward the 2020 goal (<https://etsc.eu/un-agrees-on-road-safety-sub-targets-to-aid-progress-on-2020-sustainable-development-goals/>. Accessed 20.01.2020). As noted by the Organization for Economic Co-operation and Development (OECD) and the International Transport Forum (ITF), this is the “strongest ever mandate for action on road safety” (ITF 2016, p. 17). Road safety targets have also been adopted at the national and city/state levels in Sweden (Tingvall and Haworth 1999; Johansson 2009; Belin et al. 2012; McAndrews 2013), Norway (Elvik 2008), the Netherlands (Wesemann et al. 2010), Australia (Corben et al. 2010), the United States (Cushing et al. 2016; Evenson et al. 2018; NYC 2018), the European Union (EU) (European Commission 2011; Tolón-Becerra et al. 2014), and many other countries (see Part 2 of this handbook).

The frequent use of goals and targets in road safety policy and management makes it interesting to investigate under what conditions the adopted goals and

targets can be expected to be achieved. The preconditions for successful management by objectives (MBO) in the public sector have been studied extensively. As a result, an array of factors leading to inadequate goal fulfilment have been identified, including insufficiently operational goals and targets (Elvik 1993; Edvardsson and Hansson 2005) and insufficiently calibrated goal evaluation methodologies (Larsson and Hanberger 2016), gaming behavior (Smith 1995; Propper and Wilson 2003; Bevan and Hood 2006), and lack of adequate communication channels between different governance or administrative levels responsible for implementing such goals (Wibeck et al. 2006). Much of this literature has addressed other policy areas apart from transportation and road safety, such as environmental and climate policy (Lundqvist 2004; Edvardsson Björnberg 2009, 2013) or public health in a broader sense (Lager et al. 2007; Smith and Busse 2010). However, exceptions exist. As early as 1993, Elvik addressed the question of what criteria road safety targets must satisfy in order to better guide the rational choice of means (1993). A decade later, Rosencrantz et al. (2007) used a set of criteria, which resembled the politically more established SMART criteria, to analyze the rationality of the Swedish Vision Zero for traffic safety. Drawing on the literature in the field, Elvik (2008) identified seven conditions for successful road safety MBO, some of which concerned the goals themselves (“challenging, yet in principle achievable,” “there should not be too many targets in view of the available policy instruments designed to realize them,” etc.) while others concerned the policy context in which the goals were to be implemented (“responsible agencies should be supplied with sufficient funding to implement all cost-effective road safety measures,” “incentives should exist to ensure commitment to targets from all agencies responsible for realizing them,” etc.) (Elvik 2008, p. 1116). Sobis and Okouma (2017) analyzed the use of MBO in transportation service for the disabled in the municipality of Gothenburg (Sweden). They identified two factors that contributed to its successful outcomes: that MBO was used in combination with other steering practices and the establishment of a new organizational culture that emphasized the importance of information exchange between politicians and civil servants regarding goal evaluation and achievement.

This chapter summarizes the main themes and conclusions of research that have been conducted on goal-setting in road safety policy and management to date. The primary focus of the chapter is on the rationality of the goals and targets themselves, rather than on the policy context in which the goals are to be implemented. Drawing on previously published research on rational goal-setting (Edvardsson and Hansson 2005; Rosencrantz et al. 2007; Edvardsson Björnberg 2016), the chapter aims to address the following research questions:

- What are the potential benefits of using goal-setting in road safety management?
- When is a road safety goal or target rational (functional), that is, what criteria must the goal fulfil in order to be operative?
- What is the ideal number of road safety goals within a system of goals?
- How should road safety goals and sub-goals be aligned?
- Why does a road safety goal have to be stable over time?
- Who should be involved in the road safety goal-setting process?

Before proceeding to the research questions, the next section presents a brief history of the development of road safety policy and the use of road safety goals and targets in the United Kingdom and the United States. In the following, the terms “goal,” “target,” and “objective” will be used interchangeably.

A Brief History of Goal-Setting in Road Safety Management

The invention of the modern car is commonly attributed to the German engineer and designer Karl Benz. His *Benz Patent-Motorwagen*, for which he received a patent in 1886, is considered by many to be the first power vehicle using an internal combustion engine. In the early days of motorcar invention, only limited numbers of cars for each brand were produced and used on the roads, which meant the number of car accidents and fatalities was low (Accidents involving coaches and pedestrians were obviously a cause of concern even before the motorcar was invented, thus leading to the implementation of safety measures, such as prohibiting minors from being cart drivers (Gregersen 2016). In one tragic instance, Pierre Curie, husband of Marie Curie, was killed when he fell under a horse-drawn cart in Paris in 1906.). It was not until large-scale manufacturing was introduced in the early twentieth century that road traffic situations remotely similar to the present one arose. Among the early organizations that noticed the problem of road safety were the National Safety Council (NSC) in the United States and the London Safety First Council (known today as the Royal Society for the Prevention of Accidents, RoSPA) in the United Kingdom. The NSC was established in 1913 with the aim of promoting health and safety initially in the workplace until the organization eventually expanded its efforts to include road safety, among other things. The London Safety First Council was established in 1916 in connection with a meeting convened by the operating manager of the London General Omnibus to address the “alarming increase in traffic accidents, and the direct connection therewith of the restricted street lighting which had been necessitated by the War conditions” (www.rospace.com; Jackson 1995).

Much of the work carried out by the early road safety organizations targeted the behaviors of individual road users, to whom accidents were causally attributed. In line with this, it was agreed that the behavior of individual road users, particularly children and drivers, had to be modified to fit the current road transport system. Among the early road safety measures initiated by the RoSPA were educational campaigns and road safety competitions targeting school children and professional drivers, respectively (www.rospace.com). The “traditional” behaviorally oriented approach to road safety management became very influential in the early twentieth century and is still the predominant approach to road safety management in many countries (Johansson 2009; Belin et al. 2012). However, other types of road safety measures have also been implemented, including improvements in road infrastructure and vehicle design (Oster Jr and Strong 2013). For instance, electric red-green traffic lights were introduced in the United States in 1912–1914 (https://en.wikipedia.org/wiki/Traffic_light#History. Accessed 09 Jan 2020). In 1922, the London Safety First Council suggested improvements in street lighting and argued

for the marking of road crossings frequently in use (www.rospace.com). The first seat belt patent in the United States was secured by Edward J. Claghorn in 1885, but it would take another 65 years before seat belts were made available in American-made cars. In Britain, vehicle braking requirements were introduced by the Motor Car Act 1903, and further requirements concerning the construction, weight, and equipment of cars were introduced by the Road Traffic Act 1930.

Therefore, the early road safety work was, to a significant degree, advocated for and implemented by voluntary organizations. Although national and federal state governments were involved in establishing laws on speed limits, vehicle registration, and driver licensing—notable examples were the British Locomotive Acts (1861–1878) and the New York drunk driving laws (1910)—it was not until much later, around 1930, that national and federal state governments started to organize their work with the pronounced aim to increase road safety (For more details on the specifics of these regulations, see The Locomotives on Highways Act 1861, The Locomotive Act 1865, and the Highways and Locomotives (Amendment) Act 1878. See also Motor Car Act 1903 and Road Traffic Act 1930.) A first step in this direction was taken in 1933 when the British Government decided to start investigating the causes of road accidents. In fact, the RoSPA had already put forward calls for accident causation analysis to the British Minister of Transport in as early as 1928.

In the decades that followed the Second World War, road safety work became increasingly systematic and institutionalized. Road safety policies and plans were developed, coordinated, and implemented by governmental agencies given a parliamentary or congressional mandate to improve road safety. In the United States, the Department of Transportation was created by the Congress in 1966 with the mission to “serve the United States by ensuring a fast, safe, efficient, accessible and convenient transportation system that meets our vital national interests and enhances the quality of life of the American people, today and into the future.” Four years later, in 1970, the NHTSA was established by the Highway Safety Act and was given the responsibility to reduce deaths, injuries, and economic losses resulting from motor vehicle accidents.

A crucial feature of the systematic and institutionalized approach to road safety policy was the introduction of quantified road safety targets. The idea of MBO had been developed by Peter F. Drucker in the 1950s as a way of managing business corporations (Drucker 1954) (In the academic literature, especially in the Nordic countries, the term “management by objectives and results” (MBOR) is frequently used (e.g., Lundqvist 2004; Steineke and Hedin 2008; Kristiansen 2015). “Performance management” is a related management philosophy (Ammons and Roenigk 2015).). The basic idea of MBO is that business corporations can be effectively and efficiently managed if work departures from and is evaluated against a set of objectives set by senior managers and then implemented by the employees. In the 1980s, MBO was gradually implemented in the public sector, where it subsequently became a central feature of the so-called new public management (NPM) (Hood 1991) (Kristiansen (2015) argues that the origins of MBO in the Nordic countries can be traced back to the 1960s, 1970s, or 1980s, depending on which aspect of MBO is

accentuated. In Sweden, for example, MBO-related ideas were already discussed in the 1960s in relation to budget reforms, although a comprehensive MBO reform was not launched until the 1980s (Sundström 2006).). In a public sector context, MBO typically encompasses a mixture of political control and administrative autonomy and discretion (Christensen and Laegreid 2001; Lundqvist 2004; Sundström 2006):

- Policy goals are adopted by politicians, while the responsibility for implementing them lies with the subordinate bodies. The implementing agencies enjoy a considerable degree of discretion in deciding on which measures to take to achieve the goals.
- Progress toward goal achievement is monitored, measured, and reported back to the politicians on a regular basis.
- The politicians are responsible for making policy adjustments and strategic decisions based on the reported results, for instance, on resource allocation (MBO systems typically reward agencies whose results are satisfactory and, conversely, punish agencies with insufficient goal achievement; however, this is not a pronounced feature of MBO in, for example, the Swedish public sector, where failure to achieve a policy goal is often regarded as a reason for adding funds (Steineke and Hedin 2008).

Thus, translated into a road safety policy context, MBO essentially means that road safety goals are formulated by the politicians, while the responsibility for implementation, evaluation, and feedback is delegated to another government body, in many cases the national or state road transport administration. The underlying rationale for governing by objectives rather than by rules or direct instructions is efficiency, both in economic terms and based on the belief that the administration knows how to best effectuate the intentions of the politicians.

It is obviously difficult to locate the exact point in time when a national government first quantified road safety targets. Nevertheless, the Lalonde Report, published by the Canadian Minister of National Health and Welfare in 1974, is sometimes identified as a milestone in this development (Belin et al. 2010). In the Lalonde Report, setting quantitative public health targets was introduced as a means of stimulating and coordinating governmental efforts toward reduced mortality and morbidity. The goal-setting strategy set out in the report was later taken up by the World Health Organization (WHO) and put to use by national public health authorities, which began to adopt quantified road safety targets (OECD/ITF 2016). Sweden was one of the early adopters, with quantified road safety targets being discussed as early as 1972 and eventually decided by the Parliament in 1982 (Belin et al. 2010).

Why Use Quantified Goals in Road Safety Management?

As noted above, early road safety work was not policy-driven but largely uncoordinated and implemented by different actors in an ad hoc fashion. Measures were certainly introduced with a clear goal in mind, namely, to reduce the number of

deaths and injuries caused by road accidents. However, these measures were typically not part of a strategy or plan adopted with the aim of coordinating actions across time or among agents. Nevertheless, in many respects, the road safety measures introduced in the early days of widespread automobile use were successful in saving lives and averting harm. This brings forward the question of why road safety goals and targets had to be introduced at all. More generally, what are the potential benefits of introducing goal-setting as a management tool in road safety work?

Goals are typically set because the individual or organization (henceforth “agent”) who sets the goal wants to achieve the state of affair referred to in the goal and because they believe that it becomes easier to achieve the desired states of affairs by setting the goal. Edvardsson and Hansson (2005) use the term “achievement-inducing” to denote those goals that contribute to their own achievement, i.e., goals that are rational (functional, or operative). There are two ways by which goals can be achievement-inducing. First, an adopted goal can be used to plan and coordinate action toward goal achievement over time and between agents. With the help of a goal, a road safety organization can plan and allocate work tasks to different departments or administrative units while ensuring that everyone knows what to do, when to do it, and how what is being done fits into what other departments or units are doing. For instance, a goal such as Vision Zero for road safety can be used to guide the selection of strategies or the adoption of goals and targets further down in the administrative chain (Tingvall and Haworth 1999). Moreover, an adopted goal can be used to induce and sustain action among the organization’s employees. It may invigorate the commitment of the organization and its employees to road safety. Thus, Wong et al. (2006) argue that the role of road safety targets is “to provide a basis for motivating and monitoring actions to reduce death and injury in road traffic crashes” (p. 997).

Past studies have provided considerable empirical evidence to support the idea that goal-setting can have a positive effect on an individual’s performance and, hence, be conducive to goal achievement (Locke and Latham 1990, 2002). Goals that are specific and challenging have, for instance, been shown to affect an individual employee’s choices, efforts, and persistence in such a way that goal achievement is furthered (Latham et al. 2008). Although significantly less empirical evidence exists on public policy goals and how they affect organizational output (Jung 2014), studies suggest that the positive relationship between goal-setting and goal achievement may also hold true in this case. Elvik (1993) analyzed how road safety performance differed among Norwegian counties during the years 1982–1985 and 1986–1989 based on whether quantitative or qualitative road safety targets had been adopted. He showed that quantified road safety targets were more successful than qualitative targets in reducing the accident rate per kilometer; moreover, “the best performance was achieved by counties with highly ambitious quantified targets” (p. 569). Wong et al. (2006) investigated the association between goal-setting and road safety improvement during the period 1981–1999. In their study, they investigated 14 countries that had adopted road safety management targets and examined the road fatalities before and after the setting of

the target in each country. The results showed that there was a significant overall reduction in road fatalities after quantified road safety targets had been adopted. On the basis of these findings, the authors argued that road safety goal-setting “helps to raise concern about road safety in societies, encourages decision-makers to formulate effective road safety strategies, and ensures that sufficient resources are allocated to road safety programs” (p. 1004) (The results were later updated by Allsop, Sze, and Wong (2011), who made some changes in the numerical estimates used. However, to the author’s understanding and as pointed out by the authors of the 2011 publication, the changes did not alter the main argument of the previous paper.). In line with these studies, the ITF (2008) concluded that countries that have adopted quantitative road safety targets do better than countries with no such targets.

However, governance based on goals and targets also has some potential drawbacks. Bevan and Hood (2006) identified several ways in which public MBO systems (in their case, public health service targets) can be vulnerable to gaming. Among other things, they discussed the practice, common among public agencies, to adopt goals based on what has been achieved recently. However, this practice has an unfortunate consequence: managers and politicians who expect to gain a renewed term of office tend not to exceed adopted goals even if they could, as that would put greater expectations on their future performance. It remains unclear to what extent Bevan and Hood’s (2006) findings can be translated into a road safety management context. Their analysis focused on the English public health service, which at the time of writing was managed through a system combining goals and targets with awards and punishments for (in)sufficient goal achievement. Such awards and punishments are not used in Swedish road safety work, to take one example; hence, the risk of gaming is arguably smaller. However, it is important to keep in mind that gaming could potentially occur in any MBO system depending on how it is set up, specifically on whether it grants awards and dispenses punishments based on performance outcomes.

When Is a Road Safety Target Rational?

Goals, therefore, have an important role in directing and motivating action, over time and among agents. However, to fulfil this action-guiding and action-motivating role, the goals must satisfy certain criteria. In the management and psychology literature, those criteria are often referred to as the “SMART criteria,” according to which goals should be specific, measurable, attainable, relevant, and time-bound. (See Rubin (2002) for an elaboration of what the SMART acronym stands for.) Edvardsson and Hansson (2005) proposed a set of rationality criteria that are similar to the SMART criteria. They argued that goals should be precise, evaluable, approachable, and motivational. Other criteria are also conceivable. Edvardsson Björnberg (2016) examined the extent to which it is considered rationally justified for goals to be stable over time. Rosencrantz (2008) and Edvardsson Björnberg (2009) suggested that additional criteria, such as consistency, comprehensiveness, and non-redundancy, may apply to

systems of goals. Elvik (1993) identified two sets of requirements that a target must satisfy to serve as a basis for the rational choice of means. The first concerns the relationship between the target and the values or preferences that the target expresses. The target should be “operational,” which means that the underlying preferences should satisfy the requirements of transitivity and completeness. The second set of requirements concerns the relationship between the target and the state of affairs to which the target refers. Elvik (1993) identified three such requirements: (1) the targets should not be self-contradictory, by which he means that they should not be “formulated in a way that makes their fulfillment logically impossible” (cf. Hansson et al. (2016) on self-defeating goals); (2) they should not be tautologically fulfilled, i.e., “formulated in a way that will make any outcome fulfill the target”; and (3) they should refer to “targetable” outcomes, i.e., “not to outcomes that are essentially by-products” (p. 570).

Whether or not a goal will have the capacity to guide and motivate action does not only depend on the characteristics of the goal itself, such as precision or measurability. The context in which the goal is set and the process of setting such a goal could also determine how effective it will be in regulating action. This appears to be the case not the least in a road safety management context. Elvik (2008) identified seven conditions for successful road safety MBO, some of which concern the goals and targets themselves, while others relate to the policy context in which they are to be implemented. Among the contextual factors identified by Elvik are endorsement of the goals by the top management (e.g., politicians, ministries of transport, road safety authorities), availability of funding sufficient for implementing the goals, and the establishment of a system for monitoring progress and providing feedback to the responsible agencies, among others.

Below, the two categories of rationality criteria—for individual goals and systems of goals—will be discussed separately before proceeding to the contextual conditions.

Precision

A goal can guide action only when the implementing agent knows what the state of affairs referred to by the goal is and to what extent the actions undertaken bring him/her closer to those state of affairs. This requires that the goal is both precise and possible to evaluate. Edvardsson and Hansson (2005) argued that, of these two criteria, precision is more fundamental, because evaluability presupposes that the desired state of affairs referred to by the goals is reasonably clear. The criterion of precision can be divided into at least three subcategories. Edvardsson and Hansson (2005) distinguished between directional, complete, and temporal precision, where directional and complete precision correspond to the SMART criterion of specificity and temporal precision to the criterion that a goal should be “time-bound.” Directional precision means that the goal specifies in what direction the implementing agent should go in order to reach the goal. Consider the following example:

1. There should be a decrease in the number of people killed on the road.

Assuming that the number of fatalities can be determined in a reasonably authoritative way from 1 year to another, the goal can be said to have directional precision. However, the goal does not tell the agent to what degree the goal should be achieved. Thus, it lacks complete precision. Consider instead the following goal:

2. Nobody should be killed in the road traffic system.

This goal has both directional precision (decrease in the number of fatalities) and complete precision (nobody). However, it lacks temporal precision, because it does not tell the agent within what time period the goal should be achieved. The following goal has directional, complete, and temporal precision:

3. The number of people killed on the roads should be decreased by 5% annually between 2020 and 2030.

Similar distinctions between different forms of goal precision have been made by other authors.

Elvik (1993) distinguished between “open targets,” “semi-open targets,” and “closed targets.” Open targets are qualitative and are not specified to any degree or time. They correspond to Edvardsson and Hansson’s (2005) directionally precise targets. Semi-closed targets are either quantified in time or have a quantified level component. They correspond to Edvardsson and Hansson’s (2005) temporally and completely precise targets. Finally, closed targets have both temporal and complete precision.

Imprecise goals, or “goal ambiguity,” refer to a fairly common phenomenon in public organizations (Chun and Rainey 2005; Rainey 2014). (See Chun and Rainey (2005, p. 2) for a definition of “organizational goal ambiguity.”) Carrigan (2018) identified a number of reasons behind the relative proliferation of ambiguous goals in public policy. Public policy goals are typically the result of political compromise. In political contexts, keeping a goal vague is advantageous, because one may more easily gain broad support for it than if it is specified in greater detail. From a political point of view, it may be more expedient to propose a vague goal on the basis of which political unity can be sought and then work out the details at a later point when a reasonable degree of consensus surrounding the importance of the issue has been achieved. Thus, Chun and Rainey (2005) concluded that “goal clarification is often considered ‘managerially sound’ but ‘politically irrational’ in the public sector” (p. 23). Moreover, goals can be specified as part of the implementation process. Sometimes, adopting a goal and working out the details are more efficient, especially after accumulating knowledge about the policy issue and how it can be addressed. Many implementation issues can be difficult to foresee; therefore, it may be more efficient to allow the implementing agencies to specify precisely what to achieve (cf. Lindblom 1959). Finally, politicians may find it helpful to keep goals imprecise, as policy pronouncements in the form of clearly stated goals are often used as benchmarks against which performance is measured and votes are driven.

Organizational goal ambiguity has been considered problematic in the goal-setting literature. Imprecise goals are obviously more difficult to follow up and

evaluate, and it may be more difficult for a group of agents working together toward the goals to organize their activities if the goals are imprecise and the agents lack the necessary background knowledge to coordinate their efforts efficiently. Imprecise goals may also decrease organizational commitment. Thus, Jung and Ritz (2014) argue that an ambiguous goal could “make it difficult for public employees clearly and effectively to determine how much and what kinds of effort they should make for the organization and the attainment of organizational goals, or in which direction to give effort and how to make task plans” (p. 467). (See Jung and Rainey (2011) for a discussion of the relationship between goal precision and government employee motivation.)

How precise a goal needs to be in order to have the capacity to guide action typically varies depending on the social context and the implementing agents’ background knowledge. Psychological studies have shown that for individuals, precise goals are typically more conducive to goal achievement than “do-your-best” goals; however, there are exceptions (Locke and Latham 2002). When the implementing agent is in a learning process and does not yet have sufficient knowledge to work toward an outcome-oriented goal, it may be better to set “do-your-best” outcome goals and then supplement them with specific learning goals. In those situations, specific high-performance outcome goals can be detrimental to goal achievement, as they can detract the agent from the search for an appropriate strategy. Consequently, Latham et al. (2008) concluded that “when effective behavioral routines have yet to be developed, a specific high learning goal rather than a performance one should be set” (p. 390). These findings may have some bearing on the issue of goal-setting in road safety management. Arguably, in countries and organizations with little previous knowledge about road safety management and where those responsible for implementing the targets are in a learning process, it could be more effective to adopt specific high learning goals (related to road safety management) rather than some specific performance outcome goal.

Evaluability

While precision concerns the goal (desired end-state) itself, evaluability concerns the agent’s actions and their effects on goal achievement. In this sense, precision is a more fundamental criterion, because evaluability depends on it. Goal-setting theory assumes that goals regulate performance more reliably when work is evaluated and information about how far one has come in relation to the goal is fed back to the goal-setter and/or implementer. Feedback and feedback mechanisms operate on different levels (Edvardsson and Hansson 2005). First, an agent or group of agents who are given information about where they stand in relation to the goal can more easily adjust their actions so as to further goal achievement more effectively. Second, such information can also be used to revise the goal itself. In many situations, it is difficult to know in advance whether or not a certain goal, for example, a road safety management target, is reasonably ambitious (see below). In such cases, information from evaluations may be necessary to adjust the level of goal difficulty. Third,

adopting goals that are both precise and possible to evaluate is a prerequisite for establishing accountability for insufficient goal achievement. (Hence, for the incentive for political decision-makers to adopt vague goals that are difficult to evaluate, see above.) Finally, goals that are evaluable and evaluated could have a motivating function; see below.

Successful goal evaluation presupposes both that the goal itself (the desired end state) is clear and that it is possible to assess the degrees of goal achievement. In some situations, one and the same goal can be evaluated on the basis of more than one parameter. This is, for instance, the case with the Swedish Vision Zero for road traffic safety. The Swedish Vision Zero states that nobody should be killed or seriously injured on the road. One could imagine a situation wherein there is a decrease in the number of killed people while the number of seriously injured people is increasing at the same time, or vice versa. It is not entirely clear how such a mixed result should be interpreted in terms of actual goal achievement (Rosencrantz et al. 2007).

Approachability

In goal-setting theory, it is commonly argued that goals ought to be realistic in the sense that it should be possible to at least approach them to a meaningful degree (Edvardsson and Hansson 2005). Locke and Latham (1990, 2002), for instance, provide empirical support for the so-called high performance cycle, that is, the idea that better goal achievement can be reached when the goals are precise and challenging yet not excessively difficult to achieve (see also Latham and Locke 2007). Elvik (1993) followed this line of argument when he stated that road safety management targets should be “challenging, yet in principle achievable” (p. 1116) and that “any quantified target is a compromise between idealism and realism” (p. 579).

Formulating challenging but sufficiently realistic goals is not an easy task, especially not in the public sector wherein goal achievement is dependent on many different factors. This is also the case for road safety management targets. Setting “optimally challenging” road safety targets presupposes that the goal-setter has some knowledge about what social developments affecting transportation and traffic safety can be envisaged during the period between goal-setting and projected goal achievement, what (if any) road safety management measures will likely be introduced during the indicated time period (what measures are politically and economically feasible), and how effective the introduced measures will be in furthering goal achievement (Wesemann et al. 2010; see also Corben et al. 2010).

A common argument against Vision Zero for road traffic safety is that it is an unrealistic goal, because we will never be able to achieve zero fatalities and serious injuries in road traffic (e.g., Long 2012). This argument is analyzed by Abebe et al. 2021, Edvardsson Björnberg, and Hansson (► Chap. 3, “Arguments Against Vision Zero: A Literature Review” of this handbook). As noted there, several counter-arguments could be made against this way of reasoning. It could be argued that, when it comes to saving lives and avoiding serious injuries, no goal above zero

should be considered ethically permissible. The Swedish Vision Zero for road safety policy is premised on the ethical assumption that it is unacceptable for people to be killed or seriously injured in the road transport system (Elvebakk 2007). Thus, Belin et al. (2012) argued that, “politically and humanly speaking, it was difficult to stipulate any other long-term goal” (p. 173). Trying to figure out what constitutes an optimal target level in terms of the number of killed and seriously injured people simply does not seem like a morally acceptable approach. Rosencrantz et al. (2007) drew parallels on how goals are set in workplace health and safety. Although few would deny that compromises between costs and workplace safety are unavoidable, government agencies are seldom (if ever) instructed to find out what constitutes an economically optimal level of fatal workplace accidents. Instead, it is assumed that a serious workplace accident should always be regarded as a failure.

Moreover, it is worth noting that many important political goals, such as liberty, social justice, or sustainable development, are highly idealistic in the sense that they concern end states that cannot be achieved once and for all but will have to be fought for indefinitely (Rosencrantz et al. 2007). This has not prevented political leaders and movements from formulating and using them as goals. Tingvall and Haworth (1999) appear to follow this line of reasoning when they write that zero (as in Vision Zero) is not a target to be achieved by a certain date. As noted by Kerr and LePelley (2013), highly ambitious goals, sometimes referred to as “stretch goals,” have also been adopted in business organizations in order to stimulate creativity and “out-of-the-box thinking” among the organization’s employees (see also Sitkin et al. 2011).

Motivity

A goal can be achievement-inducing not only by virtue of its action-guiding capacity; it can also facilitate goal achievement by motivating the agent to work toward the goal. The action-motivating function of goals is central to goal achievement. Merely knowing what has to be done to reach the goal is not sufficient for it to be achieved; the agent(s) responsible for working toward the goal must also want, or be motivated, to do so.

Arguably, inducement to take measures toward goal achievement could come from sources other than the goal itself (Edvardsson and Hansson 2005). For example, government agencies are often bound by legal rules prescribing that certain measures be taken, in which case the agencies and their employees have little choice but to follow the rules. Inducement to take measures that facilitate goal achievement could also arise if the organization and its employees stand to gain financially or in some other way from efforts taken to reach the goal. Moreover, leadership style and the process through which goals are adopted within an organization can have an impact on the employees’ motivation and willingness to work toward set goals (Bronkhorst et al. 2015; see below). Although these are important factors significantly affecting motivation within organizations and among employees, it is possible that such motivation could also be generated from the goal itself.

What makes a goal motivating is, of course, largely due to its content. Vision Zero for road traffic safety is motivating because many people perceive its end state to be desirable or worth striving for. However, psychological studies confirm that the motivating capacity of a goal can also come from other aspects of the goal than its end state. Locke and Latham (1990, 2002), among others, have showed that the motivating capacity of a goal is tightly linked to the other goal criteria discussed above, that is, it largely depends on how well the goal satisfies the criteria of precision, evaluability, and achievability. For instance, some studies have shown that many agents' goals that are precise and challenging exert a higher degree of motivation (both in terms of intensity and durability) than, for example, do-your-best goals (Locke and Latham 1990, 2002; Wright 2004). Studies have also presented evidence suggesting that goals that are evaluable and evaluated have a positive effect on performance, which is believed to originate from the motivational effect experienced by individuals as they are able to determine where they stand in relation to the goal (Locke and Latham 1990, 2002).

Assuming that a goal's motivating capacity is a function of how well it satisfies the other rationality criteria, it could be questioned whether motivity should really be considered an independent goal criterion. For a goal to be motivating, it would then suffice to ensure that it is sufficiently precise, evaluable, and possible to achieve to a satisfactory degree. However, some academic scholars have argued that there could be additional dimensions of goal motivity that are not adequately captured by the other criteria. Nutt and Backoff (1997), for example, discuss what makes a corporate or organizational vision motivating. Their argument is also relevant for organizational goals. In the authors' view, the degree of motivation exerted by a vision partly depends on its articulation. Here, articulation is understood in broader terms than precision. By formulating a vision through expressive images that directly "crystallize in people's mind what is wanted," a shared understanding of an organization's direction can be promoted among its employees (Nutt and Backoff 1997, p. 314).

Vision Zero for road traffic safety (i.e., nobody should be killed or seriously injured on the road) has the advantage of being both precise in the sense discussed in previous sections and formulated through an expressive image, as elaborated by Nutt and Backoff (1997). (See Rosencrantz et al. (2007) for a more extensive analysis of the precision of the Swedish Vision Zero.) The term "zero" is used here to communicate the message that it can never be ethically acceptable that people are killed or seriously injured in the road transport system (Tingvall and Haworth 1999).

To see how the criterion of motivity could encompass additional dimensions than goal precision, consider the following two goals:

1. In 2020, Sweden should have the lowest number of road traffic fatalities and serious injuries within the EU.
2. Sweden should decrease the number of road traffic fatalities and serious injuries to 15% between 2015 and 2020.

Given that adequate accident data are available, both goals could be said to be precise. However, goal 1 might in addition have the advantage of being slightly more

motivating, assuming that the competitive feature of the goal (being ranked at the top of the EU nation list) has this effect. Empirical studies are obviously needed to confirm this. However, the example serves to illustrate that there could be more to the criterion of motivity than what can be captured by the criteria of precision, evaluability, and approachability.

Balancing the Criteria

In summary, a goal can be achievement-inducing (rational, functional, or operative) by virtue of being action-guiding or action-motivating, or both. Edvardsson and Hansson (2005) argued that improved satisfaction of each of the four criteria, namely, precision, evaluability, approachability, and motivity, *ceteris paribus* makes the goal function better in the achievement-inducing sense. However, conflicts between the criteria could occur. The criteria that make a goal action-guiding could, for instance, make it less motivating and vice versa. Edvardsson and Hansson (2005) provided some examples of such conflicts. Although they will not be repeated here, the main message is that the four criteria need to be balanced from case to case in order to optimize the achievement-inducing function of the goal. Factors beyond the goal itself, such as the implementing agent's background knowledge, will determine the extent to which the four criteria need to be satisfied in each individual case to further goal achievement.

One way of balancing the action-guiding and action-motivating properties of a goal is to adopt goal systems in which some goals primarily motivate action, while others guide action. As an example, in the case of the Swedish Vision Zero, the Swedish Parliament has adopted an overarching, assumingly motivating, vision—"Vision Zero"—that is operationalized through a more precise interim target: "In 2020, no more than 220 people shall die and 4,100 people shall be seriously injured on the Swedish roads."

Rationality Criteria for Systems of Goals

Public policy goals are seldom adopted in isolation but are parts of a system of goals addressing a certain policy problem or issue. One such goal system is the Swedish system of transport policy objectives. It consists of a general transport policy objective ("To ensure the economic efficiency and long-term sustainability of transport provision for citizens and enterprise throughout Sweden") that is operationalized into a "functional objective" and an "impact objective," which, in turn, are specified through time-bound interim targets (Government Offices of Sweden 2016, p. 6) (The *functional objective* states that "The design, function and utilisation of the transport system are to provide everyone with a basic level of accessibility of good quality and usability and to contribute to the development potential of the entire country. The transport system is to be gender equal, meaning it is to meet the transport needs of women and men in an equivalent manner." The *impact objective* states that "The design, function

and utilization of the transport system are to be adapted in such a way that no one is killed or seriously injured and further the achievement of the overarching generational goal for the environment and environmental quality objectives and contribute to improving human health.”).

Another example is the UN’s SDGs, which consist of 17 broad goals and 169 targets to be achieved by 2020. For systems of goals, such as the Swedish transport policy objectives and the SDGs, additional rationality criteria to those discussed above may apply. Below, three such criteria will be discussed: completeness, number of targets, and consistency.

Completeness

In public policy goal systems, goals on a higher administrative level are often operationalized or broken down into sets of sub-goals on a lower level. The question then arises: How should the goals and targets in the system be aligned? A commonly defended idea is that the goals and targets should be aligned such that the higher goal is achieved or at least approached to a significant degree if all targets on a lower level are reached. This requires that the targets, taken together, are complete in the sense that they capture the most salient aspects of the overarching goal. If this does not hold, achieving the targets may give the impression that work toward the desired end state is progressing, while the opposite is true in reality. (See Tingvall et al. (2010) for a similar discussion of the relationship between road traffic safety performance indicators (SPIs) and the overarching goal of creating a safe road transport system.)

The adequate operationalization of an overarching goal requires that the operationalizing agents possess a good understanding of what causes the policy issue, or problem, and how it can be addressed, among other things. For example, the operationalization of the overall goal of road safety (i.e., to avoid fatalities and serious injuries in road traffic) presupposes that one knows who is involved in those accidents and what factors contribute to causing harm to the people involved. Mononen and Leviäkangas (2016) pointed out that one serious shortcoming in road safety work globally is that, although half of all deaths that occur in road traffic are suffered by vulnerable road users, mainly pedestrians and cyclists, very few road safety goals target those groups. Instead, the adopted goals emphasize in-vehicle safety, among others, which primarily affects driver and passenger safety.

When the overarching goals are qualitative, additional problems may arise in the process of goal operationalization. Qualitative goals, such as “maintaining a flourishing flora and fauna” or “providing a high-class education to everyone,” are often operationalized through precise and more easily evaluable quantitative targets. In those situations, goal displacement may occur. This happens when the goal-setting and implementing agencies lose sight of the overarching goal and instead treat the quantitative targets as if they were the “ultimate” goals (Bohte and Meier 2000). Goal displacement has at least two unfortunate implications. Unless the quantitative targets are complete in the sense that they can be said to cover the most central aspects of the qualitative goal, achievement of the subordinate targets

does not represent achievement of the overarching goal. Moreover, focusing on the quantitative aspects of a policy problem could conceal its political nature and make it appear as if the problem is strictly technical in nature (Cortner 2000). Framing a policy problem as predominantly “political” versus “technical” has consequences for whom is expected and allowed to participate in the efforts to solve the problem.

Admittedly, goal displacement may not be a serious problem in road safety policy, as many road safety targets are of a quantitative kind, even when adopted on an overarching level, as in the case of the Swedish Vision Zero. However, the problem of goal displacement could hypothetically occur if qualitative road safety goals are formulated and operationalized through quantitative sub-goals or targets.

The Number of Targets

One way of ensuring that all aspects of the overarching goal are adequately operationalized is to adopt as many targets as possible, with each target representing some aspect of the overarching goal. Thus, the requirement of completeness appears to favor goal systems consisting of a large number of sub-goals or targets. However, such goal systems have also been criticized, thus raising the question of whether there is a limit to how many goals a goal system should contain to function well in an achievement-inducing sense.

The UN’s SDGs are an example of a goal system that has been criticized for being “so sprawling and misconceived that the entire enterprise is being set to fail” (The Economist 2015). One problem with having adopted a large number of goals is that priorities have to be made somewhere along the line. However, in the absence of clear directions for how priorities should be made, there is a risk that the agencies responsible for implementing the goals will focus on less relevant policy aspects than others.

Elvik (2008) appeared to follow this line of reasoning when he argued that one must limit the number of goals in order to create an effective system of road safety management goals. Instead of trying to address all policy problems at once, one should concentrate on a few key areas that have a major impact on safety (see also Smith and Busse 2010). Here, high-quality accident data can play an important part, thereby rendering road safety policy, including the prioritization of road traffic policy goals, to be more evidence-based.

Consistency

When goals are adopted as parts of a goal system, it may be argued that the goals should ideally be consistent in the sense that they do not conflict with one another. The problem with conflicting goals is that efforts to attain one goal make it more difficult to achieve another goal. Thus, the implementing agent will have to prioritize between the goals, which can be resource consuming unless some guidance on how priorities ought to be made has been provided by the goal-setter (e.g., the

politicians). Goal conflicts, their formal definition, and the possible ways of addressing them in public policy contexts are issues that have been addressed in past works (Rosencrantz 2008; Nilsson et al. 2016; Carrigan 2018). They will not be discussed at length in this chapter. However, two observations will be briefly presented before proceeding.

First, it is hard to determine the degree of consistency required for a goal system to be rational in the achievement-inducing sense. Complete consistency may come with its own cost. Hansson (1998) warned that attempts to avoid goal conflicts altogether can lead to the adoption of goals with very low ambition levels. Such goals may be easy to co-achieve; however, their achievement will not represent any significant change in the present state of affairs. Put differently, they are not particularly meaningful.

Second, it is worth noting that goal conflicts are addressed in different ways in various policy areas. In some policy areas, compromises and adjustments between (conflicting) goals are made before the goals are set. This is, for instance, common in economic policy. In other policy areas, the goals are set first and compromises are made afterward. This is often the case in health and safety policy. Although traffic and transportation usually follow the first pattern of decision-making, road traffic safety is an exception (Rosencrantz et al. 2007). Road traffic safety work, particularly when discussed in terms of Vision Zero, follows the second pattern. Like public health issues in general, this is based on the ethical premise that it is morally unacceptable that people are killed or seriously injured on the road when mitigating measures can be taken to avoid these.

Why Does a Road Safety Target Have to Be Stable Over Time?

If goals are to fulfil their typical function of regulating action toward goal achievement, they need to have a certain stability. This holds true for both goals set by individuals and goals adopted by organizations. Frequent goal revision makes it difficult for the individual/organization to plan their activities over time. It also makes it more difficult for them to coordinate their actions with other individuals/organizations. Thus, the non-reconsideration of adopted goals appears to be the default position both in private life and in public policy.

However, it is not difficult to think of situations in which, for instance, a government agency has reason to revise one or several of its goals. For example, a road safety target could turn out to be much more (or less) difficult to achieve than initially believed and therefore in need of some modification. To take one example, Wesemann et al. (2010) referred to a set of road safety goals adopted by the Dutch government for which the level of ambition was eventually sharpened as the government found out that an unusually strong decrease in traffic fatalities occurred some years after the goals were adopted.

Baard and Edvardsson Björnberg (2015) identified two types of considerations that could give an organization (or individual) a reason to reconsider its goals: achievability- and desirability-related considerations. On the one hand, the

organization could realize that the goal's level of ambition is too high or too low (example above) and thus decide to adjust the goal accordingly (achievability-related considerations). On the other hand, an organization could realize that the value premises upon which the goal was initially set no longer hold (i.e., the end state is no longer considered valuable or is considered much less valuable than when the goal was adopted).

In relation to road safety targets, it is more likely that achievability-related considerations may come into play and justify adjustments in the targets. It would be more difficult to envision that people's preferences and values regarding loss of human life and well-being would change significantly, even over longer time periods, although other social developments leading to more severe losses of human life, such as severe pandemics, could obviously lead to politicians prioritizing other policy measures over road traffic safety.

Obviously, new data can justify the adjustments of the adopted road safety targets. However, Elvik (1993) argued that frequent goal revisions on the basis of accident history and traffic forecasts are problematic, because, over time, such a practice can lead to the present number of accidents being adopted as the target (cf. Bevan and Hood 2006, on ratchet and threshold effects). This clearly reduces the value of the goals as policy-making instruments.

Who Should Be Involved in the Goal-Setting Process?

It is sometimes argued that participatory goal-setting is conducive to goal achievement, at least in employer-employee settings. When an employee is involved in the goal-setting process, he/she will not only be more motivated to reach the goal but also have a better idea of what to do in order to reach the goal, or so it is argued (O'Connell et al. 2011). Jung and Ritz (2014, p. 468), for instance, argued that involving employees in the goal-setting process can improve their "affective organizational commitment." Various psychological mechanisms contribute to this. As noted by the authors, involvement in the goal-setting process sends a signal to the employees that "they are important, worth-while, and valued by their superiors and the organization, which in turn can enhance their self-esteem" (ibid.). Moreover, such involvement can boost the employee's trust in and identification with the organization's goals and its values.

However, mixed empirical results have been obtained in relation to the hypothesis that participatory goal-setting leads to better goal achievement. Certainly, some studies confirmed that employees who participate in the goal-setting process generally perform better than their colleagues who are simply assigned to achieve some goals. However, Latham et al. (2008) pointed out that the reason why performance levels were increased in some of those studies was that significantly more ambitious goals were adopted when participatory goal-setting processes were used. Therefore, these studies indicate that it is the level of ambition, rather than the method of goal-setting, that determines how wide the goal-outcome gap will be. To further support their argument, Latham et al. (2008) cited a study by Latham and Steele (1983),

which concluded that “when goal difficulty is held constant, there is no motivational benefit to one method of goal setting versus the other provided that the logic or rationale for an assigned goal is given” (Latham et al. 2008, p. 388). The findings suggest that, when goals are assigned, a clearly stated rationale may fulfil a similar function to that allegedly served by participatory goal-setting, namely, to foster trust, legitimacy, and commitment to the goal.

In public policy settings, goals are commonly set by one agent, in many cases a political decision-maker, and implemented by one or several other agents. In road safety policy, those other agents can be government agencies or various road safety professionals (e.g., national road administrations), local councils, business corporations (e.g., vehicle producers), and individuals. It is sometimes argued that greater acceptance of the goals will be fostered by allowing those other actors to participate in the goal-setting process. Thus, Kristianssen et al. (2018) argued that a road safety target that is “formulated from below by actors within the policy area, can enhance the internal legitimacy of the vision” (p. 266).

Corben et al. (2010) provided an example of a participatory bottom-up approach to the adoption of road safety targets. They described how the Western Australian Government and the Road Safety Council of Western Australia proceeded with the development of a new road safety strategy for the period 2008–2020. The WA Road Safety Council decided that the new road safety strategy should be developed “in a consultative and transparent way to maximize stakeholder and community acceptance” (p. 1085 f.). The development of the strategy thus involved a great degree of stakeholder consultation: over 4,000 people participated in the consultation process, which involved three phases. In phase 1, the community’s views on road safety were gathered through a number of community forums. In phase 2, community members were given the opportunity to comment on a recommended package of initiatives developed by Monash University Accident Research Centre (MUARC). This was carried out through community forums and a survey of a representative sample of the population. In phase 3, the endorsed strategy was communicated to the public, the stakeholders, and the Parliament. Using both community forums and survey samples was considered an effective way of gathering the views of those who were specifically interested in road safety issues for some reason (forum attendants) and those who were not (survey). The authors concluded that “to promote community acceptance of the Safe System philosophy and a bold, long term vision for road safety, it has been important to share the research with the community and listen to its views” (Corben et al. 2010, p. 1095).

However, involving the public in the goal-setting and implementation process, as opposed to letting the road safety experts and politicians decide, comes with its own dangers. In the case of Western Australia, analyzed by Corben et al. (2010), the consulted community had divided views on the desirability of speed limit reductions. As a result of public opposition, the WA Road Safety Council did not implement the reductions in speed limit suggested by the MUARC. Instead, a decision was made to focus on demonstration projects to illustrate the effects of speed limit reductions. According to the MUARC’s estimates, the proposed speed limit reduction would have reduced the number of killed or seriously injured people by 1,600 over the

12-year life of the strategy. Therefore, in this case, it could be argued that wider community support and greater legitimacy were gained at the potential cost of 1,600 deaths and serious injuries.

In the context of health targets, Smith and Busse (2010) acknowledged that public consultations can be useful in identifying priorities for improvement. However, they cautioned against uncritically accommodating every interest group in the process of target setting. Regardless of policy area, a wide range of considerations and interests come into play when setting a target, including the respective interests of future users, taxpayers, and users of other services. It is the role of the government to balance those (sometimes) conflicting interests and demands. When inviting stakeholders to participate in the goal-setting process, the government must bear this responsibility in mind and actively discourage potential attempts by some stakeholder groups to “hijack” the agenda in order to further their own interests.

What About Contextual Rationality Aspects?

The extent to which goals can further their own achievement depends not only on how they are formulated or how goals and sub-goals are aligned. Instead, factors that are external to a goal may also have a significant impact on how well that goal will be able to guide and motivate action toward its achievement. A number of such factors can be envisaged, and some of them will be outlined below.

A system for monitoring and evaluating progress: One of the conditions for the successful MBO in road safety policy identified by Elvik (2008) is that there should be a system in place for monitoring progress toward the targets and providing feedback to responsible agents. Thus, effective MBO presupposes not only that the goals themselves are evaluable but also that they are, in fact, evaluated on a regular basis. There must be a system of monitoring and evaluating in force that can provide feedback to those responsible for adopting the goals, and in the case of road safety, these include mostly governments at the national and local levels. Sometimes, feedback systems include incentives and disincentives (awards and punishments) that supplement the motivating effects of the goals themselves. These could be official announcements of the performance output, “naming and shaming” or other reputational measures, award of bonuses, etc. (Bevan and Hood 2006). As noted above, evaluation systems that give feedback in the form of awards and punishments may be vulnerable to gaming.

Clear communication channels: Another prerequisite for effective MBO, as pointed out by Wibeck et al. (2006), is that there should be well-developed communication channels between agents operating at various government levels. Effective communication is needed when both implementing and evaluating adopted goals, not least to avoid misunderstandings and controversies concerning the goal content, time frames, and potential goal conflicts. Arguably, this is especially crucial when the MBO system contains qualitative goals. As an example, Wibeck et al. (2006) discussed the case of the Swedish environmental quality

objective “a good built environment,” which is interpreted in different ways by various agents depending on the administrative context and the ideological perspective. Here, Vision Zero has the advantage of being formulated in quantitative terms, which could make communication regarding goal achievement easier than in many other public policy contexts. However, effective communication channels between agents at different societal levels are nevertheless important in road safety work, not the least to facilitate discussions about the effectiveness of various road safety measures.

Strong political commitment: Finally, though perhaps trivially, for road safety goals to be achieved, it is not sufficient that the goals have the capacity to guide and motivate action, that there are well-designed systems in place for the assessment and evaluation of the goals, and that effective communication channels exist between road safety agents operating at different societal levels; beyond these, there must also be a strong political commitment to the goals in the form of sufficient resource allocation and political prioritization (Elvik 2008).

Conclusions

In this chapter, it has been argued that goal-setting can be an effective management technique in road safety policy. Goal-setting theory, developed by psychologists and management theorists over the last 40 years, supports this argument. Empirical evidence from the research field of road safety confirms that the theoretical claims made in the goal-setting literature are also relevant in a road safety policy context: ambitious quantified targets can indeed reduce the number of dead and seriously injured people on the road. However, for this to hold true, a number of requirements must be satisfied. Not only must the goals be formulated in such a way that they can guide and motivate agents to act in ways that are conducive to goal achievement. In addition, goal-setting requires that there are effective evaluation systems and practices in place; that implementation measures, assessment, and evaluation outputs are communicated and discussed among the road safety agents concerned; and that adequate financial means are allocated to those responsible for implementing and evaluating the goals. In addition, governments that wish to organize their road safety work around one or several road safety targets should keep in mind that, in many public policy settings, goals and goal-setting are vulnerable to gaming. Thus, given that gaming corrupts the point of adopting and working toward goals, it is vital to counteract such behaviors at an early stage in the goal-setting process.

Cross-References

► [Arguments Against Vision Zero: A Literature Review](#)

References

- Allsop, R. E., Sze, N. N., & Wong, S. C. (2011). An update on the association between setting quantified road safety targets and road fatality reduction. *Accident Analysis and Prevention*, 43(3), 1279–1283. <https://doi.org/10.1016/j.aap.2011.01.010>.
- Ammons, D. N., & Roenigk, D. J. (2015). Performance management in local government: Is practice influenced by doctrine? *Public Performance and Management Review*, 38(3), 514–541.
- Baard, P., & Edvardsson Björnberg, K. (2015). Cautious utopias: Environmental goal-setting with long time frames. *Ethics, Policy and Environment*, 18(2), 187–201. <https://doi.org/10.1080/21550085.2015.1070487>.
- Belin, M. Å., Tillgren, P., & Vedung, E. (2010). Setting quantified road safety targets: Theory and practice in Sweden. Retrieved from <https://www.hilarispublisher.com/open-access/setting-quantified-road-safety-targets-theory-and-practice-in-sweden-2157-7420.1000101.pdf>
- Belin, M. Å., Tillgren, P., & Vedung, E. (2012). Vision Zero – A road safety policy innovation. *International Journal of Injury Control and Safety Promotion*, 19(2), 171–179. <https://doi.org/10.1080/17457300.2011.635213>.
- Bevan, G., & Hood, C. (2006). What's measured is what matters: Targets and gaming in the English public health care system. *Public Administration*, 84(3), 517–538. <https://doi.org/10.1111/j.1467-9299.2006.00600.x>.
- Bohte, J., & Meier, K. J. (2000). Goal displacement: Assessing the motivation for organizational cheating. *Public Administration Review*, 60(2), 173–182. <https://doi.org/10.1111/0033-3352.00075>.
- Bronkhorst, B., Steijn, B., & Vermeeren, B. (2015). Transformational leadership, goal setting, and work motivation: The case of a Dutch municipality. *Review of Public Personnel Administration*, 35(2), 124–145. <https://doi.org/10.1177/0734371X13515486>.
- Carrigan, C. (2018). Clarity or collaboration: Balancing competing aims in bureaucratic design. *Journal of Theoretical Politics*, 30(1), 6–44. <https://doi.org/10.1177/0951629817737859>.
- Christensen, T., & Laegreid, P. (2001). New public management – Undermining political control? In T. Christensen & P. Laegreid (Eds.), *New public management. The transformation of ideas and practice* (pp. 93–119). Aldershot: Ashgate.
- Chun, Y. H., & Rainey, H. G. (2005). Goal ambiguity in U.S. federal agencies. *Journal of Public Administration Research and Theory*, 15(1), 1–30. <https://doi.org/10.1093/jopart/mui001>.
- Corben, B. F., Logan, D. B., Fanciulli, L., Farley, R., & Cameron, I. (2010). Strengthening road safety strategy development 'Towards Zero' 2008–2020 – Western Australia's experience scientific research on road safety management SWOV workshop 16 and 17 November 2009. *Safety Science*, 48(9), 1085–1097. <https://doi.org/10.1016/j.ssci.2009.10.005>.
- Cortner, H. J. (2000). Making science relevant to environmental policy. *Environmental Science and Policy*, 3(1), 21–30. [https://doi.org/10.1016/S1462-9011\(99\)00042-8](https://doi.org/10.1016/S1462-9011(99)00042-8).
- Cushing, M., Hooshmand, J., Pomares, B., & Hotz, G. (2016). Vision Zero in the United States versus Sweden: Infrastructure improvement for cycling safety. *AJPH*, 106(12), 2178–2180. <https://doi.org/10.2105/AJPH.2016.303466>.
- Drucker, P. F. (1954). *The practice of management*. New York: Harper.
- Edvardsson Björnberg, K. (2009). What relations can hold among goals, and why does it matter? *Crítica*, 41(121), 47–66. Retrieved from <http://critica.filosoficas.unam.mx/pg/en/index.php>
- Edvardsson Björnberg, K. (2013). Rational climate mitigation goals. *Energy Policy*, 56, 285–292. <https://doi.org/10.1016/j.enpol.2012.12.057>.
- Edvardsson Björnberg, K. (2016). Setting and revising goals. In S. O. Hansson & G. Hirsch Hadorn (Eds.), *The argumentative turn in policy analysis* (pp. 171–188). New York: Springer International Publishing.
- Edvardsson, K., & Hansson, S. O. (2005). When is a goal rational? *Social Choice and Welfare*, 24(2), 343–361. <https://doi.org/10.1007/s00355-003-0309-8>.
- Elvebakk, B. (2007). Vision Zero: Remaking road safety. *Mobilities*, 2(3), 425–441. <https://doi.org/10.1080/17450100701597426>.

- Elvik, R. (1993). Quantified road safety targets: A useful tool for policy making? *Accident Analysis and Prevention*, 25(5), 569–583. [https://doi.org/10.1016/0001-4575\(93\)90009-L](https://doi.org/10.1016/0001-4575(93)90009-L).
- Elvik, R. (2008). Road safety management by objectives: A critical analysis of the Norwegian approach. *Accident Analysis and Prevention*, 40(3), 1115–1122. <https://doi.org/10.1016/j.aap.2007.12.002>.
- European Commission. (2011). *Roadmap to a single European transport area – Towards a competitive and resource efficient transport system*. COM (2011) 144 final, 28.3.2011.
- Evenson, K. R., LaJeunesse, S., & Heiny, S. (2018). Awareness of Vision Zero among United States' road safety professionals. *Injury Epidemiology*, 5(1), 21. <https://doi.org/10.1186/s40621-018-0151-1>.
- Government Offices of Sweden. (2016). *Renewed commitment to Vision Zero: Intensified efforts for transport safety in Sweden*. Stockholm: Ministry of Enterprise and Innovation. Retrieved from: <https://www.government.se/information-material/2016/09/renewed-commitment-to-vision-zero-%2D%2Dintensified-efforts-for-transport-safety-in-sweden/>
- Gregersen, N. P. (2016). *Trafiksäkerhet: Samspelet mellan människan, tekniken, trafikmiljön*. Wolters Kluwer: Stockholm.
- Hansson, S. O. (1998). Should we avoid moral dilemmas? *The Journal of Value Inquiry*, 32(3), 407–416.
- Hansson, S. O., Edvardsson Björnberg, K., & Cantwell, J. (2016). Self-defeating goals. *Dialectica*, 70(4), 491–512. <https://doi.org/10.1111/1746-8361.12161>.
- Hood, C. (1991). A public management for all seasons? *Public Administration*, 69(1), 3–19. <https://doi.org/10.1111/j.1467-9299.1991.tb00779.x>.
- International Transport Forum (ITF). (2008). *Towards zero: Ambitious road safety targets and the safe system approach*. Paris: OECD Publishing. <https://doi.org/10.1787/9789282101964-en>.
- ITF. (2016). *Zero road deaths and serious injuries: Leading a paradigm shift to a safe system*. Paris: OECD Publishing. <https://doi.org/10.1787/978928210855-en>.
- Jackson, R. H. (1995). The history of childhood accident and injury prevention in England: Background to the foundation of the Child Accident Prevention Trust. *Injury Prevention*, 1(1), 4–6. Retrieved from <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC1067529/pdf/injprev00009-0007.pdf>
- Johansson, R. (2009). Vision Zero – Implementing a policy for traffic safety. *Safety Science*, 47(6), 826–831. <https://doi.org/10.1016/j.ssci.2008.10.023>.
- Jung, C. S. (2014). Extending the theory of goal ambiguity to programs: Examining the relationship between goal ambiguity and performance. *Public Administration Review*, 74(2), 205–219. <https://doi.org/10.1111/puar.12176>.
- Jung, C. S., & Rainey, H. G. (2011). Organizational goal characteristics and public duty motivation in U.S. federal agencies. *Review of Public Personnel Administration*, 31(1), 28–47. <https://doi.org/10.1177/0734371X10394404>.
- Jung, C. S., & Ritz, A. (2014). Goal management, management reform, and affective organizational commitment in the public sector. *International Public Management Journal*, 17(4), 463–492. <https://doi.org/10.1080/10967494.2014.958801>.
- Kerr, S., & LePelley, D. (2013). Stretch goals: Risks, possibilities, and best practices. In E. A. Locke & G. P. Latham (Eds.), *New developments in goal setting and task performance* (pp. 21–31). New York: Routledge.
- Kristiansen, M. B. (2015). Management by objectives and results in the Nordic countries: Continuity and change, differences and similarities. *Public Performance and Management Review*, 38(3), 542–569. <https://doi.org/10.1080/15309576.2015.1006468>.
- Kristiansen, A. C., Andersson, R., Belin, M. Å., & Nilsen, P. (2018). Swedish Vision Zero policies for safety – A comparative policy content analysis. *Safety Science*, 103, 260–269. <https://doi.org/10.1016/j.ssci.2017.11.005>.
- Lager, A., Guldbrandsson, K., & Fossum, B. (2007). The chance of Sweden's public health targets making a difference. *Health Policy*, 80(3), 413–421. <https://doi.org/10.1016/j.healthpol.2006.05.005>.

- Larsson, M., & Hanberger, A. (2016). Evaluation in management by objectives: A critical analysis of Sweden's national environmental quality objectives system. *Evaluation*, 22(2), 190–208. <https://doi.org/10.1177/1356389016638751>.
- Latham, G. P., & Locke, E. A. (2007). New developments in and directions for goal-setting research. *European Psychologist*, 12(4), 290–300. <https://doi.org/10.1027/1016-9040.12.4.290>.
- Latham, G. P., & Steele, T. P. (1983). The motivational effects of participation versus goal setting on performance. *Academy of Management Journal*, 26(3), 406–417. <https://doi.org/10.5465/256253>.
- Latham, G. P., Borgogni, L., & Petitta, L. (2008). Goal setting and performance management in the public sector. *International Public Management Journal*, 11(4), 385–403. <https://doi.org/10.1080/10967490802491087>.
- Lindblom, C. E. (1959). The science of 'muddling through'. *Public Administration Review*, 19(2), 79–88. Retrieved from <https://www.jstor.org/stable/973677>.
- Locke, E. A., & Latham, G. P. (1990). *A theory of goal setting and task performance*. Englewood Cliffs: Prentice Hall.
- Locke, E. A., & Latham, G. P. (2002). Building a practically useful theory of goal setting and task motivation. *American Psychologist*, 57(9), 705–717. <https://doi.org/10.1037/0003-066X.57.9.705>.
- Long, R. (2012). *The zero aspiration, the maintenance of a dangerous idea*. Kambach: Human Dimensions. Retrieved from <https://safetyrisk.net/wp-content/uploads/downloads/2012/08/The-Zero-Aspiration.pdf>
- Lundqvist, L. J. (2004). Management by objectives and results: A comparison of Dutch, Swedish and EU strategies for realising sustainable development. In W. M. Lafferty (Ed.), *Governance for sustainable development: The challenge of adapting form to function* (pp. 95–127). Northampton: Edward Elgar Publishing.
- McAndrews, C. (2013). Road safety as a shared responsibility and a public problem in Swedish road safety policy. *Science, Technology & Human Values*, 38(6), 749–772. <https://doi.org/10.1177/0162243913493675>.
- Mononen, P., & Leviäkangas, P. (2016). Transport safety agency's success indicators – How well does a performance management system perform? *Transport Policy*, 45, 230–239. <https://doi.org/10.1016/j.tranpol.2015.03.015>.
- Nilsson, M., Griggs, D., & Visbeck, M. (2016). Map the interactions between Sustainable Development Goals. *Nature*, 534, 320–322. <https://doi.org/10.1038/534320a>.
- Nutt, P. C., & Backoff, R. W. (1997). Crafting vision. *Journal of Management Inquiry*, 6(4), 308–328. <https://doi.org/10.1177/105649269764007>.
- NYC (New York City). (2018). *Vision zero: Year four report*. Retrieved from <https://www1.nyc.gov/assets/visionzero/downloads/pdf/vision-zero-year-4-report.pdf>. Accessed 20 Jan 2020.
- O'Connell, D., Hickerson, K., & Pillutla, A. (2011). Organizational visioning: An integrative review. *Group and Organization Management*, 36(1), 103–125. <https://doi.org/10.1177/1059601110390999>.
- Oster, C. V., Jr., & Strong, J. S. (2013). Analyzing road safety in the United States. *Research in Transportation Economics*, 43(1), 98–111. <https://doi.org/10.1016/j.retrec.2012.12.005>.
- Propper, C., & Wilson, D. (2003). The use and usefulness of performance measures in the public sector. *Oxford Review of Economic Policy*, 19(2), 250–267. <https://doi.org/10.1093/oxrep/19.2.250>.
- Rainey, H. G. (2014). *Understanding and managing public organizations*. San Francisco: Wiley/Jossey-Bass.
- Rosencrantz, H. (2008). Properties of goal systems: Consistency, conflict, and coherence. *Studia Logica*, 89(1), 37–58. <https://doi.org/10.1007/s11225-008-9117-6>.
- Rosencrantz, H., Edvardsson, K., & Hansson, S. O. (2007). Vision Zero – Is it irrational? *Transportation Research Part A: Policy and Practice*, 41(6), 559–567. <https://doi.org/10.1016/j.tra.2006.11.002>.

- Rubin, R. S. (2002). Will the real SMART goals please stand up? *The Industrial-Organizational Psychologist*, 39(4), 26–27. Retrieved from <http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.523.6999&rep=rep1&type=pdf>
- Sitkin, S. B., See, K. E., Miller, C. C., Lawless, M. W., & Carton, A. M. (2011). The paradox of stretch goals: Organizations in pursuit of the seemingly impossible. *Academy of Management Review*, 36(3), 544–566. <https://doi.org/10.5465/amr.2008.0038>.
- Smith, P. (1995). On the unintended consequences of publishing performance data in the public sector. *International Journal of Public Administration*, 18(2–3), 277–310. <https://doi.org/10.1080/01900699508525011>.
- Smith, P. C., & Busse, R. (2010). Learning from the European experience of using targets to improve population health. *Preventing Chronic Disease*, 7(5), A102. Retrieved from <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC2938396/#!po=34.8485>
- Sobis, I., & Okouma, O. G. V. (2017). Performance management: How the Swedish administration of transportation for the disabled succeeded. A case study of transportation service for the disabled, the Municipality of Gothenburg. *NISPAcee Journal of Public Administration and Policy*, 10(1), 141–175. <https://doi.org/10.1515/nispa-2017-0007>.
- Steineke, J. M., & Hedin, S. (2008). *Management by objectives and results: Structures and practices in the regional policy field in the Scandinavian countries and Iceland*. (Nordregio working paper 2008: 3). Stockholm: Nordregio.
- Sundström, G. (2006). Management by results: Its origin and development in the case of the Swedish state. *International Public Management Journal*, 9(4), 399–427. <https://doi.org/10.1080/10967490601077178>.
- The Economist. (2015). The 169 commandments. Leader published in the 28th March 2015 issue.
- Tingvall, C., & Haworth, N. (1999). Vision Zero: An ethical approach to safety and mobility. 6th ITE International Conference on Road Safety and Traffic Enforcement: Beyond 2000, Melbourne, 6–7 September 1999.
- Tingvall, C., Stigson, H., Eriksson, L., Johansson, R., Krafft, M., & Lie, A. (2010). The properties of Safety Performance Indicators in target setting, projections and safety design of the road transport system. *Accident Analysis and Prevention*, 42, 372–376. <https://doi.org/10.1016/j.aap.2009.08.015>.
- Tolón-Becerra, A., Lastra-Bravo, X., & Flores-Parra, I. (2014). National road mortality reduction targets under European Union road safety policy: 2011–2020. *Transportation Planning and Technology*, 37(3), 264–286. <https://doi.org/10.1080/03081060.2013.875277>.
- United Nations. (2010). *Improving global road safety*. Resolution adopted by the General Assembly on 2 March 2010. A/Res/64/255. <https://undocs.org/A/RES/64/255>. Accessed 20 Jan 2020.
- Wesemann, P., van Norden, Y., & Stipdonk, H. (2010). An outlook on Dutch road safety in 2020; future developments of exposure, crashes and policy. *Safety Science*, 48(9), 1098–1105. <https://doi.org/10.1016/j.ssci.2010.05.003>.
- WHO (2018). Road safety. Retrieved from <http://www.who.int/features/factfiles/roadsafety/en/>. Accessed 20 Jan 2020.
- Wibeck, V., Johansson, M., Larsson, A., & Öberg, G. (2006). Communicative aspects of environmental management by objectives: Examples from the Swedish context. *Environmental Management*, 37(4), 461–469. <https://doi.org/10.1007/s00267-004-0386-1>.
- Wong, S. C., Sze, N. N., Yip, H. F., Loo, B. P. Y., Hung, W. T., & Lo, H. K. (2006). Association between setting quantified road safety targets and road fatality reduction. *Accident Analysis and Prevention*, 38(5), 997–1005. <https://doi.org/10.1016/j.aap.2006.04.003>.
- Wright, B. E. (2004). The role of work context in work motivation: A public sector application of goal and social cognitive theories. *Journal of Public Administration Research and Theory*, 14(1), 59–78. <https://doi.org/10.1093/jopart/muh004>.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.



Zero Visions and Other Safety Principles

2

Sven Ove Hansson

Contents

Introduction	33
The “Zero Family”	33
Zero Defects	34
Workplace Safety	35
Traffic Safety	38
Crime Prevention	40
Preventive Medicine	43
Environmental Protection	47
Disarmament	49
Comparisons	50
Improvement Principles	55
Continuous Improvement	55
As Low as Reasonably Achievable (ALARA)	57
Best Available Technology (BAT)	60
Summary	62
Aspiration Principles	63
Risk Limits	63
Exposure Limits	64
Process and Equipment Regulations	66
Cost-Benefit Analysis	67
Individual Cost-Benefit Analysis	70
Cost-Effectiveness Analysis	70
Hypothetical Retrospection	71
Summary	72

S. O. Hansson (✉)

Division of Philosophy, Royal Institute of Technology (KTH), Stockholm, Sweden

Department of Learning, Informatics, Management and Ethics, Karolinska Institutet,
Stockholm, Sweden

e-mail: soh@kth.se

© The Author(s) 2023

K. Edvardsson Björnberg et al. (eds.), *The Vision Zero Handbook*,
https://doi.org/10.1007/978-3-030-76505-7_2

Error Tolerance Principles	72
Fail-Safety	73
Inherent Safety	75
The Substitution Principle	77
Safety Factors	79
Multiple Safety Barriers	82
Redundancy	84
Summary	85
Evidence Evaluation Principles	85
The Precautionary Principle	86
A Reversed Burden of Proof	88
Risk Neutrality	89
“Sound Science”	90
Summary	91
Conclusion	91
Cross-References	92
References	92

Abstract

Safety management is largely based on safety principles, which are simple guidelines intended to guide safety work. This chapter provides a typology and systematic overview of safety principles and an analysis of how they relate to Vision Zero. Three major categories of safety principles are investigated. The *aspiration principles* tell us what level of safety or risk reduction we should aim at or aspire to. Important examples are Vision Zero, continuous improvement, ALARA (as low as reasonably achievable), BAT (best available technology), cost-benefit analysis, cost-effectiveness analysis, risk limits, and exposure limits. The *error tolerance principles* are based on the insight that accidents and mistakes will happen, however much we try to avoid them. We therefore have to minimize the negative effects of failures and unexpected disturbances. Safety principles telling us how to do this include fail-safety, inherent safety, substitution, multiple safety barriers, redundancy, and safety factors. Finally, *evidence evaluation principles* provide guidance on how to evaluate uncertain evidence. Major such principles are the precautionary principle, a reversed burden of proof, and risk neutrality.

Keywords

ALARA · As low as reasonably achievable · Aspiration principles · BAT · Best available technology · Burden of proof · Continuous improvement · Cost-benefit analysis · Cost-effectiveness analysis · Error tolerance principles · Evidence evaluation principles · Exposure limits · Fail-safety · Improvement principles · Inherent safety · Multiple safety barriers · Precautionary principle · Redundancy · Reversed burden of proof · Risk limits · Risk neutrality · Safety barriers · Safety factors · Safety principles · Substitution · Vision zero

Introduction

Much safety work is based on *safety principles*, which are usually simple rules or mottos such as “inherent safety,” “fail-safe” and “best available technology.” Many such principles have been proposed, and there is a considerable overlap between them (Möller et al. 2018). In this chapter, we will see Vision Zero as one of the safety principles and relate it to other such principles. We will begin with its closest relatives and then consider some of its more distant kin. The safety principles discussed in this chapter are summarized in Fig. 1.

The “Zero Family”

The idea that there should be nothing at all of something undesirable must have been close at hand to human thinking since long before recorded history. Concepts of “naught” and “none” are much older than the mathematical concept of zero. For instance, at least since Alcidamas (fourth century BCE), abolitionists have claimed that no single human being should be a slave.

In modern discussions on safety, strivings to get completely rid of something undesirable have emerged in many contexts, probably often independently. The goal of “zero” is therefore an oftentimes reinvented wheel. As Gerard Zwetsloot and his co-authors have pointed out, we have a “family of zero visions” (Zwetsloot et al. 2013, p. 46). Its members are known under a great variety of names, some of which are recorded in Fig. 2. Some of them are “visions” in the sense of being difficult or perhaps impossible to fully achieve. Others are clearly possible to achieve and might

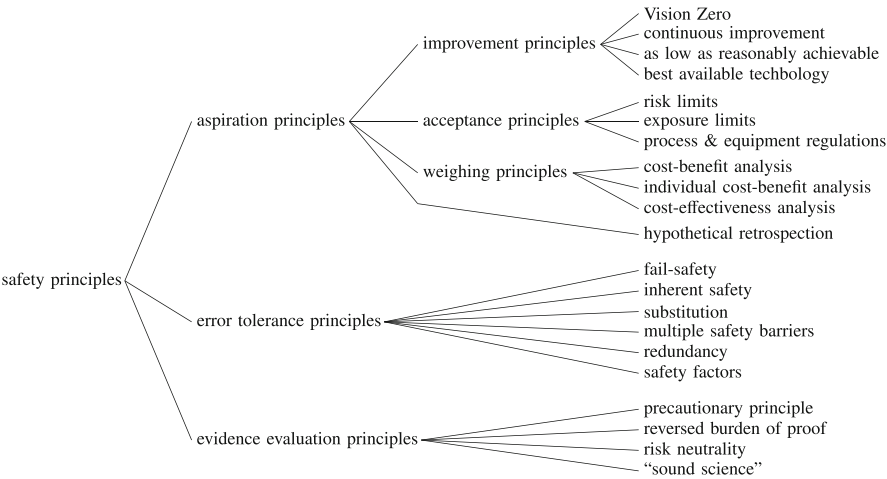
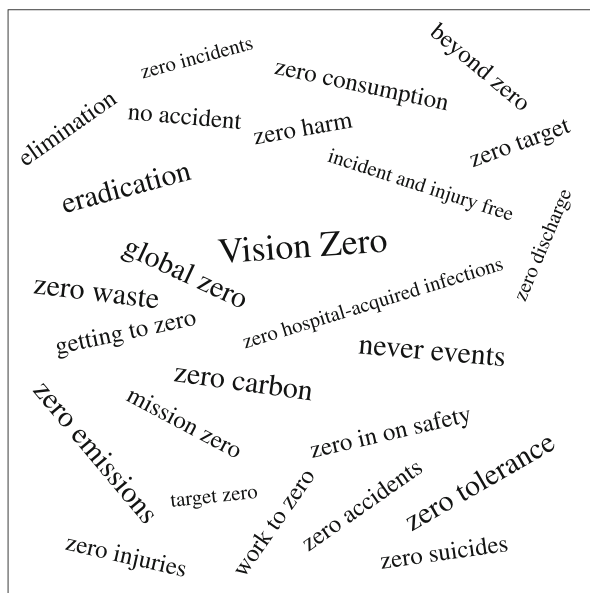


Fig. 1 A typology of the major safety principles discussed in this chapter

Fig. 2 Some of the major members of the “zero family” of safety principles



more appositely be called “zero targets.” Both visions and targets can serve useful social purposes. Visions can inspire us to undertake long-term projects whose end points have to be further specified as we go along. Targets, on the other hand, are essential parts of our planning for what to do next (Edvardsson and Hansson 2005; ► Chap. 1, “Vision Zero and Other Road Safety Targets”).

Zero Defects

The safety-related zero concepts have been influenced by zero visions in the area of industrial quality. In 1965, the US Department of Defense issued a handbook to be used by defence contractors for “establishing and implementing Zero Defects” as a “motivational approach to the elimination of human error.” According to the handbook, the programme had been originated in 1962 by a major (unnamed) defence contractor, which had “established goals for each department to reduce to zero those defects attributable to human error.” The programme had been successful, and its ideas had already been “adopted by numerous industrial and Department of Defense activities” (Anon 1965, p. 3). Each individual worker was asked to “accept voluntarily a challenge to do an errorless job,” a challenge that would be accepted by those feeling “pride in workmanship” (ibid.). Targets and scorekeeping were essential components of the concept.

The unnamed defence contractor was probably Martin Marietta, a Florida-based company that built missiles for the armed forces. Their employee James F. Halpin has been identified as the inventor of the concept (Peierls 1967). He published a book (Halpin 1966) that made the Zero Defects concept known in wider circles. Halpin

was critical of what he called a “double standard” in our attitudes as consumers and as workers. Most consumers expect the products they buy to be free of defects, but the same persons, in their role as workers, consider a certain amount of errors in their work outputs to be acceptable. Zero Defects would remove that contradiction and at the same time improve the quality of working life by making workers more proud of their work (Zwetsloot et al. 2013, pp. 45–46).

Another important promoter of Zero Defects was Philip B. Crosby, who worked in the same missile company as Halpin in the early 1960s. In his book *Quality is Free* (Crosby 1979), he emphasized the hidden costs of low quality and argued that a Zero Defects strategy is beneficial for business. The Zero Defects programme was much in vogue in American industry in the late 1960s and early 1970s, but after that, it receded in importance. Largely due to Crosby’s influence, it had a renaissance in the automobile industry in the 1990s (Lovrencic and Gomišček 2014, p. 4). Claims have been made that the “zero” concepts of safety engineering have developed out of Zero Defects (Butler 2017, p. 25), but no evidence seems to have been presented that confirms this genealogy of the concepts.

Workplace Safety

A large number of zero concepts have been used to express the goal that no one should be injured or harmed in their workplace. Common terms are “zero harm,” “zero injuries,” “zero accidents,” “zero incidents” and “incident and injury free” (abbreviated IIF) (Zou 2010; Lovrencic and Gomišček 2014). If the phrases are interpreted literally, there is a large difference, for instance, between “zero accidents” and “zero incidents,” but in actual practice, the choice among these phrases does not say much about how safety work is conducted in the workplace. The variations in terminology seem to mark different national and industry-based traditions rather than different approaches to safety.

In the **United States**, a two-year safety programme named “Zero in on Safety” was launched in 1971 in all federal agencies. Its aim was to reduce the number of injuries and other losses on federal workplaces (Anon. 1971). However, the expression “zero in on” seems to have been used in the common sense “concentrate attention on,” rather than referring to a zero target.

The American construction industry has been prominent in promoting zero approaches to workplace safety. In 1989, the Construction Industry Institute (CII) started a project called “Zero Accidents Techniques,” which led up to a “zero injury” target that was launched in 1993. It was based on detailed analyses of accidents on construction sites and means to prevent them. The programme had an emphasis on worker compliance. For instance, it included a drug-testing policy according to which a worker with a positive drug test was expelled from the workplace in the following 60 days (Hinze and Wilson 2000; Lovrencic and Gomišček 2014). In an interview, the CEO of Bechtel, one of the largest construction companies in the United States, explained the rationale behind the Zero Accidents objective as follows:

I sincerely believe all accidents and all injuries are preventable. Accidents don't just happen. They occur primarily because of someone's unsafe behaviour. Correcting that behaviour is the only way we'll get to Zero Accidents. Zero Accidents means exactly that – zero. When it comes to preventing accidents, nothing less than perfection will do. (Zou 2010, p. 15)

However, this company also has a more empowering approach to employees' contributions to safety. It has authorized all employees to stop work which they consider to be unsafe; "If it's not safe, don't do it" (Zou 2010, p. 14). Emmitt J. Nelson, who was chair of the original CII Zero Accidents Task Force, has explained the justification of zero targets with the following hypothetical example:

Last year, a small 100-employee company experienced five serious lost-workday cases. Accident costs were high, but the misery that followed the injuries was even more devastating. The owner vowed to take action.

After consulting with company leaders, the owner set a goal of only two lost-workday cases for the upcoming year. At the next safety meeting, the owner voiced his concern about the five lost-workday cases and announced the new goal. Everyone was enthusiastic and seemed to buy into the plan.

Question: On the first workday in January, how many of the 100 employees think (as far as the goal is concerned) that it is okay for them to have an injury that results in a lost-workday case?

Answer: All 100. Each employee thinks that, provided the goal of two is not reached, it is acceptable for a serious injury to occur.

What has the new goal inadvertently accomplished? It has said it is acceptable for injuries to occur (provided no more than two occur). This certainly is not the 'Let's stop injury' message the owner intended to deliver. . .

Zero is the correct approach. Such commitment sends an unmistakable message to all employees that injury is unacceptable. (Nelson 1996)

In 2019, the National Safety Council (NSC) launched a campaign called *Work to Zero 2050*. Its purpose is to achieve zero fatalities on workplaces by the year 2050. The NSC pointed out that when the campaign was introduced, the number of workplace fatalities was 5000 per year, as compared to 50,000 per year in the early 1900s, although the number of people working had increased from 30 million to 160 million in the same period. This was largely due to technology innovation, and new technology will be an essential component of the new campaign. "Its purpose is to *eliminate* death on the job by the year 2050. Period. No hedging, no qualifiers" (McElhattan 2019).

In **Japan**, a government-sponsored Zero-Accident Total Participation Campaign (Zero-Accident Campaign) was introduced in 1973 by the Japan Industrial Safety and Health Association (JISHA). Reportedly, the campaign was inspired by the American "Zero in on Safety" campaign. It was also influenced by the quality improvement movement (JICOSJ n.d.). (Dekker (2017, p. 125) claims that the Japanese "Zero-Accident Total Participation Campaign" took place in the 1960s, but this is not corroborated by the Japanese sources I have had access to.) The target of the campaign was set high:

'Zero accidents' means to achieve an accident free workplace (not only no fatal accidents or accidents causing absent from work, but also no accidents, including industrial accidents,

occupational illness, and traffic labor accidents) by detecting, understanding, and solving all hazards (problems) in everybody's daily life as well as potential hazards existing in workplaces and work. (JICOSJ [n.d.](#))

In the **United Kingdom**, the first recorded use of a zero goal for workplace safety seems to have been in 1988. In that year, the British Steel plant in Teesside in North East England launched a programme of "total quality performance." The training material used in meetings with the whole workforce included a short text on accident prevention with the two headings "Total Quality is no Accident" and "Zero Defect = Zero Accidents" (Procter et al. [1990](#)). Thus, safety was included as part of a campaign whose primary goal was product quality. The campaign was based on the notion of "continuous improvement" (see subsection "Continuous Improvement"). The site manager had visited an American steel plant that used the slogan "Total Quality is no Accident," which was the inspiration for including safety in the quality campaign. The safety team in Teesside adopted a "zero-accident philosophy" for their work. A considerable reduction in accidents was reportedly achieved (Ball and Procter [1994](#)).

In later years, the British construction industry has taken a leading role in applying zero approaches to safety. The best-known example is the construction of the venues of the 2012 Summer Olympics and Paralympics in London. The Olympic Delivery Authority adopted a "zero tolerance" approach to unsafe and unhealthy working conditions on their building sites. This was a five-year project involving 12,000 workers who worked a total of 80 million hours. The accident rate was unusually low for a building project, and there was no fatality. This seems to have been the first fatality-free Olympic construction project in modern history. The Olympic Delivery Authority received a special reward from the Royal Society for the Prevention of Accidents for this achievement (Wright [2012](#)).

Most major building companies in the United Kingdom have adopted a zero aim for their safety work. A variety of two-word brands are in vogue, such as "zero harm," "mission zero," "target zero," "zero target" and "beyond zero" (Sherratt [2014](#)). The last-mentioned catchphrase may be somewhat surprising, given the considerable difficulties that the building industry has had in eliminating accidents leading to fatalities and severe injuries. Fred Sherratt, a researcher in construction management, noted that the Beyond Zero programme "boldly announces 'Zero incidents? We can do better than that!' on its webpage." This, she says

...could suggest achieving zero incidents is an easy target. This reading was identifiable elsewhere in the text, 'aiming for zero accidents was a soft target and was not the final word in what could be achieved'. Here, Zero is arguably belittled beyond itself: positioned as 'soft', something so easily attainable that it should not be considered a target, just something to be looked beyond. When considered in the context of one of the highest risk industries in the UK... this appears to be rather empty rhetoric. (Sherratt [2014](#), p. 742)

In **New Zealand**, the New Zealand Aluminium Smelters Limited (NZAS) introduced zero target thinking in 1990 and adopted the slogan "Our Goal is Zero." This resulted in a considerably lower accident rate than that of most other aluminium smelters in the world (Young [2014](#)).

In **Australia**, most zero-aiming safety activities are performed under the designation of “zero harm,” which was introduced in an influential agreement in 2002 between the Australian Chamber of Commerce and Industry (ACCI) and the Australian Council of Trade Unions (ACTU). The term “zero harm” has become so popular that it has to a large extent replaced references to safety. Workplace health and safety personnel are recruited to positions as “zero harm manager,” “zero harm reporting coordinator” and “zero harm advisor.” Reporting of accidents and incidents is done on a “zero harm reporting app,” and safety culture is referred to as “zero harm culture” (Butler 2017, p. 28). In a PhD thesis on the Zero Harm movement in Australia, Keith Butler noted that “Zero Harm has been popularised as a mantra,” but he recognized that “industry is also actively implementing Zero Harm as a goal or vision and even as a numerical target based on the concept that if a single day without an injury can be achieved, then 365 days without an injury is also achievable” (ibid., p. 2).

In **Finland**, the Finnish Zero Accident Forum was established in 2003 as a voluntary network to help workplaces promote health and safety. It has now changed name to the Finnish Vision Zero Forum. Its membership consists of 440 workplaces, whose 450,000 employees comprise 16% of the country’s workforce (<https://www.ttl.fi/en/vision-zero-not-a-numerical-goal-but-a-mindset/>. Accessed August 19, 2020). Members of the Forum have improved their safety performance, as measured in terms of lost-time accidents, whereas non-members were on average less successful in this respect (Zwetsloot et al. 2017b, p. 22).

In **Sweden**, the government introduced a Vision Zero strategy for workplaces in 2016. Its focus is on fatal accidents, and the central formulation is as follows:

No one should have to die as a result of their job. Concrete measures are necessary in order to prevent work-related accidents leading to injury or death. (Quoted in Kristianssen et al. 2018, p. 265)

The strategy requires measures to prevent work-related injuries and to improve psychosocial work environments.

Sweden also has a Vision Zero for fire safety since 2010, stating “No one shall die or be seriously injured due to fire.” Contrary to the goal for workplace safety, the fire safety goal is combined with interim goals. (Kristianssen et al. 2018, p. 263).

Traffic Safety

This subsection is included as a brief introduction to a topic that is treated in much more detail in other chapters of this handbook.

Vision Zero as a goal for traffic safety was first adopted by the Swedish Parliament in 1997. The Bill stated that “the long-term goal of traffic safety is that nobody shall be killed or seriously injured as a consequence of traffic accidents” and that “the design and function of the transport system shall be adapted accordingly” (Government Bill 1996/97, p. 137). Vision Zero has its focus on severe accidents,

i.e. accidents leading to fatalities or serious injuries. Its basic message is that as long as serious accidents still occur, there is a need to improve traffic safety. This has been described as a radically different approach from previously dominating road safety policies, in which a certain death toll in traffic was more or less openly accepted as a price for the advantages of mobility. However, it is also stated in the Bill that Vision Zero is not intended to eliminate every traffic accident that gives rise to property damages or light personal injuries (Belin et al. 2012, p. 171). This is a matter of priorities. As long as there are a large number of deaths and serious injuries in road traffic, they have to be the prime target in traffic safety work.

Vision Zero has had significant impact on the traffic safety work of the Swedish Road Administration. The following four changes on Swedish roads have resulted from systematic work to implement Vision Zero:

- *More roundabouts*: Roundabouts have become more common in intersections, in particular within population centers. Roundabouts significantly reduce vehicle velocities. If collisions take place, their consequences will be less severe than in regular intersections, due both to reduced speed and different angles of collision.
- *Roads with midrails*: The so-called 2+1 road is a three-lane road with two lanes in one direction and one in the other. A mid barrier separates traffic going in opposite directions. The direction of the middle lane alternates so that overtaking is always possible within a few kilometres. In this way, head-on collisions are prevented, which has led to a significant reduction in fatalities and serious injuries. The 2+1 road was introduced in 1998 on a route that had previously been the scene of many fatal accidents. There was much initial scepticism towards the new road design, but it has proved effective against accidents, and this road design is now widely used in Sweden. It reduces the number of fatalities by around 80% (Johansson 2009, p. 826).
- *Lower speed limits within population centers*: In order to implement Vision Zero, local municipalities have been authorized to lower speed limits to 30 km/h. The purpose of this is to reduce fatalities among unprotected road users.
- *Safer roadsides*: Efforts have been made to mitigate accidents where vehicles drive off the road. Rails have been set up, and roadsides have been cleared of dangerous objects such as boulders and trees (Vägverket 2004).

The measures taken to implement Vision Zero have led to a considerable reduction in road accidents in Sweden. The number of road traffic fatalities per year was reduced from 541 to 221 in the period from 1997 to 2019 (<https://www.transportstyrelsen.se/sv/vagtrafik/statistik/olycksstatistik/statistik-over-vagtrafikolyckor/>. Last accessed August 11, 2020).

In September 2000, the Norwegian Parliament adopted a vision of zero killed or seriously injured. Denmark adopted a similar vision with the slogan “every accident is one too many” (Færdselssikkerhedskommissionen 2000). Other countries in Europe have adopted their own variants of Vision Zero, and so have Australia and several states and major cities in the United States (Mendoza et al. 2017). Several car

manufacturers have also taken up Vision Zero as a goal for technological developments, aiming at “zero crash cars” (Zwetsloot et al. 2017c, p. 96).

Crime Prevention

The application of zero goals to the prevention of criminal and deviant behavior has mainly taken place in the United States, and it has almost invariably been associated with the phrase “zero tolerance.” The application of this phrase to anti-crime policies dates back to 1983, when forty submarine sailors were reassigned by the US Navy for having used drugs. The term was picked up by a district attorney in San Diego, who developed a programme called “zero tolerance” in 1986. The main purpose of that programme was to prosecute all drug offenders however minor their offence was. Sea vessels carrying any amount of drugs were to be impounded (Skiba and Peterson 1999, p. 373; Skiba 2014, p. 28; Stahl 2016).

The programme received the attention of members of the US government, including the Customs Commissioner William von Raab, who decided to implement a similar approach on a national level. A government committee, called the White House Conference for a Drug-Free America (WHCDFA), issued a report in June 1988, concluding: “The U.S. national policy must be zero tolerance for illegal drugs” (Newburn and Jones 2007, pp. 223–224). Initially, the programme resonated well with sentiments in large parts of the general public. However, its implementation in the day-to-day activities of customs officials was far from frictionless. In 1990, two research vessels were seized due to small amounts of marijuana that had been found on-board. This was generally recognized as disproportionate, and after that, the US Customs Service discretely discontinued its zero tolerance programme (Skiba and Peterson 1999, p. 373). But in spite of the practical problems, the political appeal of “zero tolerance” was unscathed, and the concept had already started to proliferate in other social areas.

Two campaigns focusing on violence against women, one in Canada and one in Scotland, took up the “zero tolerance” motto. In 1993, the Canadian Panel on Violence Against Women, which had been appointed by the prime minister two years earlier, presented an action plan declaring zero tolerance. By this was meant that “no level of violence is acceptable, and women’s safety and equality are priorities.” Inspired by the Canadian initiative, the Women’s Committee of the Edinburgh City Council initiated a Zero Tolerance Campaign (ZTC) in late 1992. The campaign launched posters and cinema adverts with a prominently featured Z symbol, emphasizing that male violence against women and children should never be tolerated. A Zero Tolerance Charitable Trust was established in 1995 (Newburn and Jones 2007, pp. 224–225). It is still active, waging campaigns throughout Scotland with the goal “a world free of men’s violence against women” (<https://www.zerotolerance.org.uk/about/>. Last accessed August 19, 2020).

In 1990, when the US Customs Service abandoned their zero tolerance policies, school districts across the United States were busy introducing their own versions of zero tolerance. Already in the previous year, school districts in California, New York

and Kentucky had adopted zero tolerance procedures that targeted drugs and weapons on school premises. Dissemination was rapid, and in 1993, such procedures were implemented in schools throughout the country. The scope of the policies was significantly extended after the Columbine school shooting in April 1999 in Littleton, Colorado. Throughout the country, school security was strengthened with measures such as metal detectors, increased surveillance and greater presence of security personnel. The lists of behaviors punished with school suspensions and exclusions became much longer. In many schools, misdemeanours such as swearing, truancy and dress code violations would lead to suspension (Skiba and Peterson 1999; Stahl 2016).

Consequently, the number of suspended and expelled students increased, in some cases dramatically. Media started to report on suspensions that appeared to be inordinately harsh. One twelve-year-old student was suspended for violating her school's drug policy by sharing her inhaler with a student who had an asthma attack on a bus (Skiba and Peterson 1999, p. 375). A ten-year-old girl found a small knife in her lunchbox, where her mother had placed it for cutting an apple. She realized that it was a forbidden object and immediately handed it over to her teacher, but she was nevertheless suspended for bringing a weapon to school (APA Taskforce 2008, p. 852). A five-year-old bringing a plastic toy axe to school was suspended for the same reason. Worst of all, an eleven-year-old boy died because the school's drug regulations forbade him to bring his inhaler to school (Martinez 2009, p. 155).

Research does not confirm the assumption that suspending students from school improves their law abidance and keeps them away from crime. To the contrary, numerous studies have shown school suspensions to be associated with higher risks of school dropout, failure to graduate and criminal activity (Martinez 2009, p. 155; APA Task Force 2008; Stahl 2016). In addition, zero tolerance policies in schools have turned out to be highly discriminatory. African American students have been suspended 3–4 times more frequently than other students (Hoffman 2014, p. 71; Lacoe and Steinberg 2018, p. 209). The widely held assumption that Black students earn their higher rate of school suspensions by their own behavior is not borne out by research. Instead, research shows that African American students are disciplined more than other students for less serious misbehavior (APA Task Force 2008; Skiba 2014, p. 30). However, in recent years, many states have reduced the punitive elements of their zero tolerance programmes and replaced suspensions by policies and interventions that keep the students in the classrooms (Lacoe and Steinberg 2018, pp. 207–208).

In 1991, zero tolerance policies were adopted by the federal Department of Housing and Urban Development. A very low bar was set for evicting a tenant from public housing programmes. Tenants could even lose their apartment without doing anything reproachable themselves; it was sufficient to have visitors engaging in criminal activity in or near the apartment. For instance, an elderly couple was thrown out of their home because their grandson had smoked marijuana in the parking lot. The immediate consequence of zero tolerance was described as follows by a researcher:

Creating homelessness is the [sic] one of the main outcomes of “zero-tolerance” policies. People are evicted from their housing units under these policies and they are left without any other place to live. That is, in fact, the main purpose of these policies. And since people are most often in government-supported housing programs because they do not have other options, eviction typically results in making the person homeless. (Marston 2016)

Often, those evicted had to move to another neighborhood, with negative effects on their social networks. Frequent moves are particularly problematic for families with children, who run increased risks of lower school achievement, school dropout and substance abuse (Marston 2016). In housing, just as in schools, zero tolerance policies have largely been counterproductive, fuelling rather than reducing social exclusion and criminality.

The most well-know application of zero tolerance policies took place in a number of police forces, most notably the New York Police Department. Zero tolerance policing was inspired not only by other zero tolerance programmes but also by the so-called “broken windows” approach to police work, which was introduced in the early 1980s by James Q. Wilson and George L. Kelling. They argued that “if a window in a building is broken and is left unrepaired, all the rest of the windows will soon be broken.” In the same way, they said, minor disturbances of order in a neighborhood could, if not curbed in time, lead to an uncontrollable escalation of criminality. In such a process, “[t]he unchecked panhandler is, in effect, the first broken window.” To prevent such negative developments, they proposed that police departments should cease assigning resources “on the basis of crime rates (meaning that marginally threatened areas are often stripped so that police can investigate crimes in areas where the situation is hopeless).” Instead, priority should be given to “neighborhoods at the tipping point – where the public order is deteriorating but not unreclaimable” (Wilson and Kelling 1982).

These ideas became a central part of the “zero tolerance” policies that were introduced under William Bratton, who was appointed commissioner of the NYPD in 1994. The programme had a strong focus on aggressive police crackdowns on various forms of minor misconduct in the public space, such as drunkenness, urination, squeegeeing, fare dodging and begging. A similar programme was run by the London Metropolitan Police in the King’s Cross area in December 1996. In London, police raids against minor public order offences led to the removal of beggars and inebriated and homeless people from public areas (Innes 1999; Newburn and Jones 2007).

Homicides were substantially reduced in New York City in the period 1991–1997. This decrease has often been attributed to the zero tolerance and broken window policies. However, in this period, many other changes took place that could have influenced homicide statistics. Criminological research does not confirm the usual “success story” for zero tolerance (Bowling 1999). An interesting comparison can be made with San Diego, where a more community-oriented policing programme was carried out in about the same period. The two cities experienced similar reductions in severe crime. However, whereas there was a dramatic increase in legal actions taken against the NYPD for police misconduct, no such effect occurred in

San Diego (Greene 1999). It should also be observed that resource-demanding measures against petty crimes will necessarily divert resources that might have been directed at more serious crimes. A major assumption behind zero tolerance programmes appears to be that harsh and legalistic action against juveniles committing minor offences will deter them from a criminal career. This is not substantiated by the criminological evidence. To the contrary, harsh and legalistic treatment of young offenders is associated with a larger number of rearrests and a higher future participation in crime (Klein 1986; Innes 1999; Petrosino et al. 2014).

Against this background, it should be no surprise that the use of zero tolerance concepts has decreased considerably in police work (Wein 2013, p. 4). As an example of this, the NYPD has retreated from its zero tolerance strategy for policing (Anon. 2017).

Preventive Medicine

The prevention of diseases is one of the areas where zero targets naturally spring to mind. Why should the prevention of a disease have a less ambitious goal than its complete eradication? Ideas about eradicating infectious diseases are much older than the zero concepts discussed in the previous subsections. In later years, disease eradication has also been promoted for various iatrogenic diseases. Furthermore, a zero vision much inspired by the Vision Zero of traffic safety has been introduced in suicide prevention. This subsection is devoted to these three areas of zero-aiming medical goal setting.

Apparently, the first proposal to eradicate a disease can be found in a book published in 1793 by the English physician John Haygarth (1740–1827). He proposed that smallpox could be exterminated through a combination of obligatory variolation and strict measures to prevent contagion (Haygarth 1793). Variolation consisted in infecting a person with scabs or fluid from the skin bumps of a person with smallpox. This usually resulted in a less severe disease than after natural contagion. Variolation was far from risk-free, but the vast majority survived, and they acquired immunity against the disease. Three years after Haygarth published his book, Edward Jenner (1749–1823) made his first experiment with vaccination. He infected human subjects with pus from cowpox blisters (containing what we now know to be live viruses). Cowpox is a much milder disease than smallpox, but it gives rise to immunity also against smallpox. In 1801, Jenner published a short pamphlet on vaccination in which he made the bold prediction that it was “too manifest to admit of controversy, that the annihilation of the Small Pox, the most dreadful scourge of the human species, must be the final result of this practice” (Jenner 1801, p. 8; Fenner 1993).

In modern terminology, the word eradication is used for a “permanent reduction to zero” of the worldwide incidence of an infectious disease, such that no further interventions are needed to prevent new cases. Thus, eradication is a global concept. For regional or national absence of a disease, the word elimination is used instead (Hinman 2018; Kretsinger et al. 2017).

Smallpox was indeed possible to eradicate. It satisfies a crucial condition for this, namely, that it has no animal reservoir. If there is a species of wild animals that serve as alternative hosts for a human pathogen, then that pathogen can always survive in the wild, whatever measures we take to prevent its spread among humans. But the smallpox virus can only survive in humans. Therefore, if its dissemination in the human population could be stopped through vaccination and other efficient measures, then the virus would die out. Attempts to engage the WHO in a programme to eradicate smallpox were made in 1953, but the idea was dismissed as unrealistic. It was only in 1958 that a decision was made to introduce such a programme. An internationally funded and coordinated programme was not in place until 1967. In that year, the yearly death toll of smallpox was about two million people, and the disease was endemic (regularly present) in 32 countries. The programme had three major components. The first was mass vaccination to ensure that at least 80% of the population was vaccinated in all countries where the disease was endemic. In this way, the disease would be reduced to a level where all outbreaks could be effectively contained. The second component was an efficient reporting and response organization in all countries where the disease occurred. Places with outbreaks were visited by a team that searched for additional cases and vaccinated everyone who could have been infected. The third component was an international exchange of reports that kept all participants in the programme informed of developments throughout the world (Fenner 1993; Henderson 2011).

The programme was successful. The last case of smallpox occurred in 1978, and in 1980, the world was officially declared free of smallpox. This was 179 years after Jenner's prediction that smallpox would eventually be "annihilated" by vaccination. The eradication of smallpox put an end to an immense amount of suffering that had haunted humanity since ancient times. Only in the twentieth century, at least 300 million people died from smallpox (Henderson 2011, p. D8).

In 1988, a new disease eradication programme was started, the Global Polio Eradication Initiative (GPEI). In its more than three decades of operation, it has radically reduced the number of polio cases. In 1988, there were around 350,000 cases, distributed over 125 countries. In 2018, there were 33 cases, all of which occurred in Pakistan and Afghanistan (Polio Global Eradication Initiative 2019). The campaign is now described as an endgame, but it operates under grave difficulties caused by militant anti-vaccination propaganda and violent groups targeting vaccination workers (Kaufmann and Feldbaum 2009). In the years 2012–2015, 68 government employees working with the administration of polio vaccine were killed in Pakistan (Kakalia and Karrar 2016). In 2019, the number of polio cases increased again (<https://www.who.int/news-room/detail/03-10-2019-statement-of-the-twenty-second-ihf-emergency-committee-regarding-the-international-spread-of-poliovirus>. Last downloaded August 19, 2020).

In spite of these remaining difficulties, discussions are ongoing on how the extensive infrastructure and capabilities created by the Global Polio Eradication Initiative can be used after polio has finally been defeated. The most common answer is that these resources should be retained and utilized in efforts to eradicate measles and rubella (Kretsinger et al. 2017; Cochi 2017). Measles is a major cause of child fatalities. In 2018, more than 140,000 measles deaths were reported globally,

most of them in children (<https://www.who.int/news-room/fact-sheets/detail/measles>. Downloaded Dec 7, 2019). In industrialized countries, around 3 of 1000 children who catch measles will die from the disease, but in countries with widespread malnutrition (in particular, vitamin A deficiency) and insufficient healthcare resources, the death toll can be in the range between 100 and 300 per 1000 children with the disease (Perry and Halsey 2004). Rubella is transferred from the pregnant woman to the foetus, and it is a major cause of congenital diseases. About 100,000 children are born each year with rubella-induced inborn diseases, such as deafness, severe heart diseases, glaucoma and diabetes (Lambert et al. 2015; Banatvala and Brown 2004). Both measles and rubella are clearly possible to eradicate. Neither of them has an animal reservoir, both are preventable with two doses of vaccine and both have easily detectable clinical symptoms. Major efforts are ongoing to eliminate these diseases in most regions of the world, but progress has been slow, largely due to disinformation spread by anti-vaccination propagandists (Benecke and DeYoung 2019).

In 1986, the Carter Center in cooperation with the WHO started a campaign to eradicate the Guinea-worm disease (GWD, also called dracunculiasis). This is a painful but seldom deadly disease affecting people in Africa and Asia. The infection is spread by drinking water containing water fleas that are infected by guinea worm larvae. The disease can be prevented with relatively simple measures such as providing safe water, blocking the use of infected water sources and boiling or filtering water before drinking (Tayeh et al. 2017). The eradication programme has succeeded in drastically reducing the number of cases. In 1986, 3.5 million cases were recorded worldwide, whereas the number of cases was only about 30 in the years 2017 and 2018 (<https://www.cartercenter.org/news/pr/guinea-worm-world-wide-cases-jan2019.html>. Last accessed August 19, 2020). This shows that the disease can be contained on a very low level. However, the discovery that the disease has animal reservoirs in both dogs and frogs has dampened hopes of its eradicating (Callaway 2016; Eberhard et al. 2016).

In 2011, the Joint United Nations Programme on HIV and AIDS (UNAIDS) adopted a vision of three zeros: zero new HIV infections, zero discrimination and zero AIDS-related deaths. The World AIDS Day 2011 had the motto “getting to zero” (Garg and Singh 2013). Currently, this seems to be extremely difficult to achieve (https://www.unaids.org/sites/default/files/media_asset/UNAIDS_FactSheet_en.pdf. Last accessed August 19, 2020). There is still no vaccine suitable for mass vaccination. A study of the potential for getting “close to zero” concluded that a radical reduction in the number of new cases would nevertheless be possible through “the global implementation of a bundle of prevention strategies that are known to be efficient, including anti-retroviral therapy to all who need it (which currently does not happen due to insufficient healthcare resources) and condom promotion and distribution (which is currently prevented by ruthless religious hypocrisy)” (Stover et al. 2014). As long as these hurdles remain, it does not seem possible to defeat this disease.

One might expect the eradication of severe diseases to be an unusually uncontroversial undertaking, but that has not been the case. Anti-vaccination propagandists cause considerable problems for vaccination campaigns worldwide. Their activities have delayed the eradication of polio and led to outbreaks of measles

in countries that had for long been spared from that disease (Kakalia and Karrarb 2016; Hussain et al. 2018). Recently, a group of biologists have questioned the eradication of pathogens and hosts transmitting them to humans, appealing to the ethical standpoint that “each species may have a right to exist, independent of its value to human being [sic]” (Hochkirch et al. 2018, p. 2). Their main example was the campaign to eradicate trypanosomiasis, a deadly disease caused by a parasitic protozoan spread by bites by tsetse flies. They wanted to “stimulate discussions on the value of species and whether full eradication of a pathogen or vector is justified at all” (ibid., p. 1). This discussion should be seen in the perspective that about 10% of the earth’s about six million insect species are threatened by extinction (Diaz et al. 2019). The role of disease prevention and eradication in species extinction is minuscule in comparison to other causes of reduced biodiversity.

Zero targets have also become popular in connection with iatrogenic diseases. In particular, infections spread in hospitals and clinics have been targeted in initiatives with names such as “zero hospital-acquired infections,” “zero healthcare-associated infections,” “zero tolerance for healthcare-associated infections” and “zero tolerance to shunt infections” (Warye and Granato 2009; Warye and Murphy 2008; Choksey and Malik 2004). Such zero targets have the purpose to “set the goal of elimination rather than remain comfortable when local or national averages or benchmarks are met” (Warye and Murphy 2008). In the discussion on these goals, phrases reminding of discussions on zero goals for workplace and traffic safety are common, for instance, “even one H[healthcare] A[ssociated] I[nfection] should feel like too many” (Warye and Murphy 2008). However, as evidenced by frequent use of the term “zero tolerance,” the focus on individual compliance is often more pronounced in the medical context:

Lapses in [surgical] theatre discipline were not tolerated, and this attitude was inculcated into all present; we term this ‘zero tolerance’. (Choksey and Malik 2004)

A major impediment to achieving H[healthcare] A[ssociated] I[nfection] zero tolerance has been a lack of accountability of hospital administrators and clinicians (including unit/ward/service directors). Where in the world would we be allowed to walk into the operating room and do surgery without complying with rules/regulations/culture of the operating room (e.g., strict hand hygiene, gowns, gloves, masks, sterile techniques, etc.)? Virtually nowhere. Yet, almost everywhere, I[ntensive] C[are] U[nit] directors, ward attendants, etc., commonly witness clinicians throughout their units/wards or healthcare facility (e.g., ICUs, wards, outpatient services, emergency departments, etc.) fail to comply with recommended infection control precautions and yet they say nothing. We must engage our hospital administrators and transition from a culture of benchmarking (i.e., are we as good as others like us) to a culture of zero tolerance (i.e., are we preventing all the H[healthcare]A[ssociated]I[nfection]s we can prevent). Furthermore, hospital administrators must make it clear that unit/service/ward directors will be held accountable for the HAIs that occur in their patients in their units. . . No excuses should be tolerated. (Jarvis 2007, p. 8)

There has also been criticism against the use of zero goals in healthcare, in particular concerning iatrogenic infections. No medical intervention is completely free of risk, and sometimes, an intervention is justified although the risk of infection cannot be eliminated (Worth and McLaws 2012). A focus on zero targets

may, according to some authors, be dangerous since it makes it “increasingly difficult to educate the public about the sources of risk of healthcare interventions” (Carlet et al. 2009).

The National Quality Forum in the United States, an umbrella organization of private and public healthcare organizations, conducts a campaign against “never events.” By this is meant three types of surgical mistakes: wrong site, wrong procedure and wrong person. Wrong site operations are usually either left/right mistakes or operations on an incorrect level, e.g. surgery on the wrong vertebra or (in dentistry) wrong tooth extraction. Wrong procedure operations take place on the correct site but with the wrong type of surgery. An example would be the removal of a patient’s ovaries along with the uterus, when the purpose of the operation was only to remove the uterus. Wrong patient operations depend on confusions of patients, for instance, patients with the same name. With careful preoperative and operative procedures, the risk of these never events can be substantially reduced (Michaels et al. 2007; Ensaldo-Carrasco et al. 2018). The British National Health Service (NHS) employs a more extensive list of never events, which includes, for instance, retention of a foreign object in the patient’s body after surgery, overdoses of insulin and transfusion with incompatible blood (https://improvement.nhs.uk/documents/2899/Never_Events_list_2018_FINAL_v7.pdf. Last accessed August 19, 2020).

A Vision Zero for suicide was adopted by the Swedish government in 2007. It states, “No one should find him- or herself in such an exposed situation that the only perceivable way out is suicide. The government’s vision is that no one should have to end their life” (Kristianssen et al. 2018, p. 264). In 2011, the US National Action Alliance for Suicide Prevention (NAASP) published an ambitious action plan against suicides. It set the goal Zero Suicide, by which is meant that no suicide should take place among patients under treatment in the healthcare system. It is thus less ambitious than the Swedish goal, which also covers suicides by persons who are not patients. By adopting Zero Suicide, the organization tried to achieve “a transformation of a mindset of resigned acceptance of suicide into a mindset of active prevention of suicide as an outcome of treatment. Instead of asking how not to have more suicides than usual, a Zero Suicide organization challenges itself to have no suicides at all” (Mokkenstorm et al. 2018). Both the Swedish and the American zero suicide goals have been subject to criticism by authors who consider these goals to be unrealistic (Holm and Sahlin 2009; Smith et al. 2015). For an in-depth discussion on the goal of zero suicides and strategies to approach it, see Wasserman et al. (► Chap. 37, “Vision Zero in Suicide Prevention and Suicide Preventive Methods”).

Environmental Protection

The term “zero waste” was used in environmental discussions already in the 1970s. In 1975, two American researchers described how a water purification plant could achieve “‘zero’ waste discharge,” by which they meant total recycling of all wastes generated in the plant (Wang and Yang 1975, p. 67). The phrase “zero waste” is now

quite common, often in combination with other terms to denote an area or activity that is free from waste: “zero waste campus,” “zero waste community,” “zero waste city,” “zero waste living” (Zaman 2015), “zero waste product” (Zaman 2014, 2015) and even “zero waste humanity” (Zwier et al. 2015).

However, there are widely divergent opinions on what should be meant by “zero waste.” This can be illustrated by the various ways in which the term is used about the treatment of household waste in cities. One common definition of zero waste in that context is “diversion from landfill”; in other words, that no waste goes to landfill (Zaman 2014, p. 407). As Atiq Zaman has pointed out, this is an unambitious goal since it “does not place enough emphasis on how waste can be reused as a material resource (as opposed to being incinerated, for instance)” (ibid., p. 408).

Another, much more ambitious, interpretation identifies zero waste with total recycling. On that interpretation, “[a] 100% recycling of municipal solid waste should be mandatory to achieve zero waste city objectives” (Zaman and Lehmann 2011, p. 86). Contrary to the less ambitious goal of avoiding landfill, the goal of total recycling cannot be achieved by cities and municipalities alone. Only if the products discarded by city dwellers consist of 100% recyclable material can the city achieve zero waste in this sense. Therefore, product design has to be a crucial component of the strategy (ibid., p. 86).

However, not even 100% recycling means that nature is unscathed by the city’s activities. Recycling does not necessarily mean that a product gives rise to raw material of the same quality and quantity as the material it was made from. Some authors have required that instead of becoming waste, materials should enter an “endless scheme” and “pass through the process of usefulness without losing their capacity to feed the system again after being used” (Orecchini 2007, p. 245):

The challenge is to achieve completely closed cycles. Anything but a closed cycle, which starts from useful resources and returns to them after their use, is unable to realize truly sustainable development: diffused, shared, and ideally endless for the entire human society. (Orecchini 2007, p. 246)

Material included in such cycles will not be consumed in the usual sense of the word, and therefore, systems based on these principles have been called systems of “zero consumption” (ibid., p. 245).

Zero waste and related concepts have also been applied to industrial processes. In that case as well, there are large variations in how the terms are defined. Sometimes, remarkably lenient interpretations have been used. For instance, consider the following definition of “zero discharge”:

[A] Z[ero] D[ischarge] system is most commonly defined as one from which no water effluent stream is discharged by the processing site. All the wastewater after secondary or tertiary treatment is converted to a solid waste by evaporation processes, such as brine concentration followed by crystallization or drying. The solid waste may then be landfilled. (Das 2005, pp. 225–226)

This means that a factory is said to have “zero discharge” if all its waste is put in a landfill, even if there is toxic leakage from that landfill. As this example shows, claims that an activity produces “zero” environmental harm can be severely misleading if the “zero” does not cover all environmental detriments associated with that activity.

The term “zero emissions” was used in a legal text adopted in California in 1990, namely, the Zero Emission Vehicle Mandate. It specified steps that the automobile industry had to take towards the introduction of zero emission vehicles (Kemp 2005). The definition of the term “zero emissions” is equally problematic as that of “zero waste.” For instance, battery electric vehicles are typically called “zero emissions” vehicles because they have zero tailpipe emissions of greenhouse gases. However, the production of these cars gives rise to greenhouse gas emissions, and the electricity used to charge the batteries may not have an emission-free source (Ma et al. 2012). Similarly, the claim that a biofuel is “zero-carbon” or has “net zero emissions” needs some qualification. The greenhouse effect of a carbon dioxide molecule from a biofuel is the same as that of a carbon dioxide molecule coming from coal. The advantages of the biofuel will only materialize through replanting resulting in photosynthesis that “compensates” the emissions from burning the fuel. This effect is delayed and depends on the future use of the harvested land (Sterman et al. 2018). Therefore, although replacing fossil fuel with biofuel is a clear advantage from the viewpoint of climate change mitigation, claims of net zero emissions cannot usually be validated.

Disarmament

The pacifist’s credo is a “zero”: no wars! In arms control negotiations, total abolishing of certain types of weapons has had a prominent role. In the world’s first treaty on chemical weapons, signed in Strasbourg in 1675, France and the Holy Roman Empire agreed not to use any poisoned bullets. The Geneva Protocol that went into force in 1929 prohibits all uses of chemical and biological weapons in war (Coleman 2005). This prohibition is still a cornerstone in the international law of war.

After World War II, considerable efforts have been made to outlaw and eliminate the third major type of weapons of mass destruction, namely, nuclear weapons. The very first resolution adopted by the United Nations General Assembly, on a session in London on January 24, 1946, mandated work that would lead to “the elimination from national armaments of atomic weapons and of all other major weapons adaptable to mass destruction” (United Nations 1946). A large number of attempts have been made to reinvigorate this mission. Perhaps most remarkably, at the Reykjavík Summit in 1986, Ronald Reagan and Mikhail Gorbachev agreed to work for an agreement to eliminate all nuclear weapons. However, no such agreement materialized. One of the more prominent independent initiatives towards nuclear disarmament is the “Global Zero” initiative, which was launched in Paris

in 2008 and is still highly active. It aims for a structured destruction of all nuclear weapons, ultimately resulting in a world without such weapons. In a speech in Prague on April 5, 2009, President Barack Obama expressed “America’s commitment to seek the peace and security of a world without nuclear weapons” (Holloway 2011). At the time of writing, chances of nuclear disarmament seem to be at a low point.

However, in 1999, the Ottawa Treaty against Anti-Personnel Landmines came into force. Its signatories comprise the vast majority of the world’s countries (but as yet neither China, Russia, nor the United States). Each state that has signed the treaty is committed to “never under any circumstances” use anti-personnel mines and to destroy all such mines in its possession (<http://www.icbl.org/media/604037/treatyenglish.pdf>. Last accessed August 19, 2020).

Comparisons

As we have seen, the “zero family” is broad and diverse. In this subsection, we are going to consider three major ways in which the members of this family differ from each other: (1) how realistic they are; (2) the objects of the zero goals or, in other words, what it is that one strives to make zero; and (3) the subjects, i.e. the persons or organizations tasked with achieving or approaching zero.

Realism

The zero targets we have examined above cover a wide range of degrees of realism, from the proven (but initially doubted) realism of eradicating smallpox to goals such as zero deaths on all American workplaces that seem to be exceedingly difficult to fully achieve.

Criticism of “zero” as too unrealistic is one of the recurrent themes in the literature on zero targets (► Chap. 3, “Arguments Against Vision Zero: A Literature Review”). For instance, Goh and Xie claim that the “zero defects” goal is impossible to reach due to “one fundamental characteristic of nature, namely that all natural elements are subject to variation”:

By deduction, therefore, ‘zero defect’ by itself is a pseudo-target, attractive and even seductive when brandished at management seminars, but misleading or self-deluding on the shop floor. For example, has anyone ever claimed, directly or indirectly, to have run a printed circuit board soldering machine ‘right the first time’ and obtained ‘zero defect’ in the soldered joints all the time? (Goh and Xie 1994, p. 5)

In consequence, they propose that “do it better each time” is a better slogan than “zero defects” (Goh and Xie 1994, p. 5). Similar criticism has been waged, for instance, against Vision Zero for traffic safety (Elvik 1999) and against zero targets for iatrogenic infections (Carlet et al. 2009).

Defenders of zero targets have pointed out that what initially seems to be unrealistic may become realistic if old ways of thinking are broken up. Setting

zero goals can be an efficient way to overcome fatalism (Zwetsloot et al. 2013, pp. 44–45) and to defeat “the subtle message that fatal injuries will occur and are acceptable” (Nelson 1996, p. 23). For instance, the zero suicide goal can serve to induce “a transformation of a mindset of resigned acceptance of suicide into a mindset of active prevention of suicide as an outcome of treatment” (Mokkenstorm et al. 2018, p. 750).

Obviously, partial achievement of a zero goal can be a great step forward. For instance, the drastic reduction in polio cases that has been achieved in the programme for polio eradication is already an outstanding accomplishment in terms of human health and welfare, even though the disease has not (yet) been fully eradicated. Even in cases when full attainment is not within reach, zero goals can be efficient means to inspire important improvements. In other words, goals can be achievement inducing even if they cannot be fully attained (Edvardsson and Hansson 2005; ► Chap. 1, “Vision Zero and Other Road Safety Targets”).

Our traditions and conventions concerning the realism of goals differ considerably between social areas (see Fig. 3). In some areas, the tradition is to set goals only after carefully investigating and taking into account what is feasible and what compromises with other objectives are necessary. We can call this *restricted goal setting*. For instance, goals for economic policies are usually set in this way. Another example is the setting of occupational and environmental exposure limits, which is usually preceded by careful investigations of what exposure levels can be achieved in practice.

In other areas, the tradition is instead to set goals without first determining what is in practice feasible. We can call this *aspirational goal setting*. For instance, it is desirable that no one should be exposed to violence, but unfortunately, this is not a realistic goal that can be fully achieved. However, law enforcement agencies do not

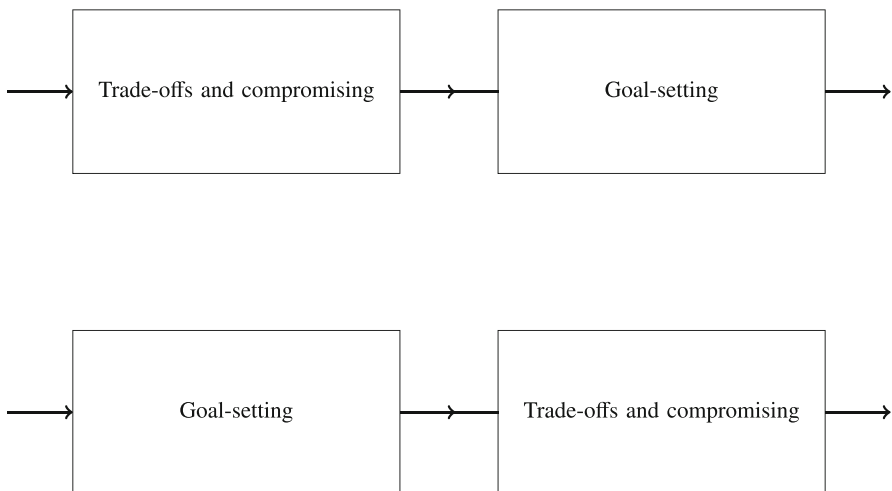


Fig. 3 Above, restricted goal setting. Below, aspirational goal setting

operate with “compromised” goals such as “at most 10 murders and 20 rapes in this district next year.” The reason for this is of course that the “uncompromised” goal of no violent crimes is good enough as an indication of what they should strive for.

It is reasonable to assume that aspirational goal setting tends to result in goals that are more inspiring than those emerging from restricted goal setting. On the other hand, the latter goals are often more suitable for guiding policy implementation and evaluation. In many cases, the best solution can be to combine both types of goals, in order to obtain both inspiration and practical guidance (Edvardsson and Hansson 2005; ► Chap. 1, “[Vision Zero and Other Road Safety Targets](#)”).

The Zero Object

By the object of a zero goal or target, we mean that which is required to be zero. Zero objects can be classified according to how narrow or broad they are. For instance, the Vision Zero of traffic safety has a fairly narrow zero object: The goal is explicitly limited to fatalities and serious injuries, and it does not require strivings for zero occurrence of less serious accidents. As we saw above, some of the measures taken to implement Vision Zero do in fact increase the frequency of less serious accidents. For instance, the introduction of roundabouts in four-way crossings decreases the risk of high-speed collisions with fatal outcomes, but it also increases the risk of low-speed collisions with at most minor personal injuries. The “Work to Zero” goal of the US National Safety Council has a similar approach; its zero object is limited to fatal accidents at workplaces (McElhattan 2019).

However, many other zero goals in workplace safety operate with a broad zero object. For instance, Emmitt J. Nelson recommends “first setting the injury commitment to zero lost-workday cases.” When that has been achieved, “the commitment can then become zero recordables” (Nelson 1996, p. 23). Even wider zero objects that include incidents (near accidents) are also common in workplace safety (Zou 2010).

Does a broad zero object, which includes incidents and minor accidents, divert attention and resources away from the most serious accidents, which should have top priority? Or is a broad approach, to the contrary, a superior way to prevent serious accidents, since it stops the beginnings of event chains that would later lead to a serious accident? The answer to these questions will depend on the extent to which the most serious accidents begin in the same way as the less serious ones, or as Sidney Dekker expressed it, “If preventing small things is going to prevent big things, then small and big things need to have the same causes” (Dekker 2017, p. 127).

In this respect, there are large differences between different types of accidents. Fires in households are an example with a large overlap. Most big fires in housing areas begin as small fires. The prevention of small home fires is therefore an efficient – and necessary – means to prevent big fires in housing areas. On the other hand, large accidents in a nuclear plant (such as those in Chernobyl in 1986 and Fukushima in 2011) do not typically start in the same way as small accidents in the same plants. Preventing accidents with handheld tools, or falls from scaffolding, in a nuclear plant, is of course important for its own sake. However, it does not usually contribute much

to preventing the event sequences that end up in large accidents with massive emissions of radioactive material to the surroundings. There are similar situations in other industries. The BP Deepwater Horizon explosion in April 2010, which killed eleven workers and gave rise to the largest marine oil spill in human history, seems to be a case in point. The accident was reportedly preceded by six years of injury-free and incident-free operations on the rig. The explosion was the result of other types of failures than those that lead to smaller workplace accidents (Dekker 2017, p. 125). Zero tolerance in law enforcement is an interesting parallel case. Its efficiency depends on the degree to which the targeted minor offences are parts of the causal chains that lead to more serious, violent crimes. The failure of some zero tolerance programmes seems to have been due to lack of such overlaps in causal chains.

We can learn from examples like these that the choice between broad and narrow zero objects has to be informed by a careful analysis of overlaps between different kinds of undesirable event chains. This may of course result in broad zero goals being chosen in some areas and narrow goals in others.

Some organizations have a whole package of zero goals, referring, for instance, to different aspects of product quality, waste reduction, etc. Zero goals for workplace safety have often been integrated as one of several zeros in such combinations. It would not be unreasonable to worry that the safety-related zero will be outcompeted by other, perhaps more economically important, goals in the same package. However, experience reported in the literature indicates that inclusion in such comprehensive zero packages strengthens the safety goal by increasing management commitment (Twaalfhoven and Kortleven 2016, p. 65; Zwetsloot et al. 2017c, pp. 95–96):

For example, Paul O'Neill, the former chief executive officer of aluminum manufacturer Alcoa, made a public commitment to his employees, the media, and investors that Alcoa would target zero accidents in its plants. He then made it the single most important performance metric for everyone in his organization, including himself. At the time O'Neill introduced this effort, many were concerned that Alcoa might be "overinvesting" in safety such that, on the margin, the benefits of added safety would be exceeded by the costs of the productivity lost in trying to achieve it. The history of Alcoa's efforts, however, suggested that there was no such trade-off between safety and productivity. In fact, O'Neill's aspirational push led to improvements on both dimensions for Alcoa. (Huckman and Raman 2015, p. 1811)

But obviously, such a conflict-free relation between safety and productivity cannot be taken for granted. Equally obviously, safety goals may have to be pursued even when they clash with other important goals. (Zwetsloot and co-workers (2017a, p. 261) report that "commitment to product safety in the Chinese industry may decrease the commitment to work safety," but no confirmation of that claim could be found in the source referred to.)

The Zero Subject

Since the early twentieth century, two major approaches to responsibilities for accidents and other untoward events have been prominent in safety management.

One of them is the *environmental theory* that received much of its scientific basis in the work that the American sociologist Crystal Eastman (1881–1928) presented in a seminal book on workplace accidents, published in 1910. Based on detailed investigations of a large number of accidents, she showed that the same types of accidents were repeated again and again and that they could be prevented by appropriate measures on the workplace. This approach put the responsibility for workplace safety on employers, in contrast with two attitudes that were common at the time, namely, that accidents were unavoidable “acts of God” and that they were the results of workers’ carelessness. Eastman’s work had considerable influence in safety engineering, and it was instrumental in the creation of a worker’s compensation law. Safety work in this tradition, with its emphasis on technological and organizational measures that make workplaces less dangerous, has been instrumental in reducing risks of accidents on all kinds of workplaces (Swuste et al. 2010).

In the 1920s, psychologists Eric Farmer and Karl Marbe independently developed another approach, often called the *accident-proneness theory*. Their basic idea was that most accidents are caused by a minority of workers who behave dangerously on the workplace. Therefore, accidents could be avoided by psychological testing that would identify and exclude these workers (Swuste et al. 2010, 2014). However, contrary to the environmental theory, the accident-proneness theory did not contribute much to improved safety. No tests were developed that could identify workers with an increased proneness to accidents. In addition, the basic assumptions of the theory turned out not to hold. For instance, much of its alleged scientific support came from studies showing that accidents are quite unevenly distributed among workers in the same workplace. This could of course depend on some workers having more dangerous tasks and working conditions than others. Such differences were not taken into account in these studies, and therefore, no conclusions can be drawn from them (Rodgers and Blanchard 1993). The theory fell into disrepute among safety professionals, but it has not entirely died out. Various attempts have been made to revive it, but it is still fraught with the problems that caused its decline in the 1950s and 1960s. In 1975, the road safety researcher Colin Cameron summarized the situation as follows:

For the past 20 years or so, reviewers have concluded, without exception, that individual susceptibility to accidents varies to some degree, but that attempts to reduce accident frequency by eliminating from risk those who have a high susceptibility are unlikely to be effective. The reduction which can be achieved in this way represents only a small fraction of the total. Attempts to design safer man–machine systems are likely to be of considerably more value. (Cameron 1975, p. 49)

This conclusion still stands. Or as Sidney Dekker, another safety researcher, concluded in a recent review of the literature:

The safety literature has long accepted that systemic interventions have better or larger or more sustainable safety effects than the excision of individuals from a particular practice. (Dekker 2019, p. 79)

The contrast between the environmental theory and the accident-proneness theory comes out clearly in the above review of zero targets and goals. Most of the zero goals have been operationalized in the tradition of the environmental theory. That applies, for instance, to Vision Zero in traffic safety and to most of the zero targets for workplace safety. However, there are also zero goals with a strong focus on correcting or removing individuals whose behavior is deemed undesirable. The clearest example of this is zero tolerance in law enforcement. Notably, applications of Vision Zero in traffic and workplace safety have on the whole been successful, whereas zero tolerance in law enforcement has largely been abandoned since it did not deliver. Against the background of previous experiences with the two major approaches to safety, this should be no surprise.

But obviously, an environmental, or as we can also call it, system-changing, approach to safety should not be taken as a reason for individuals to be careless about safety and leave everything to the system. Individual attention to safety is still needed. Furthermore, improvements of safety do not come by themselves. They require individuals who propose, demand and implement them.

Improvement Principles

The various zero targets and goals refer to different categories of risks, but their common message is that no level of risk above zero is fully satisfactory. Consequently, improvements in safety should always be striven for as long as they are at all possible. Several other safety principles have essentially the same message. Among the most prominent of these are continuous improvement, as low as reasonably achievable (ALARA) and best available technology (BAT). These principles differ in their origins, and they are used in different social areas, but they all urge us to improve safety whenever we can do so. We can call them *improvement principles* (Hansson 2019).

Continuous Improvement

The so-called quality improvement movement in industrial management originated in the United States in the 1920s and 1930s. Walter A. Shewhart (1891–1967), W. Edwards Deming (1900–1993) and Joseph Moses Juran (1904–2008) were among its most important pioneers. Their focus was on the quality of industrial products and production processes, which they succeeded in improving by means of new methods of quality control and new ways to incentivize workers. Their ideas were adopted not only in the United States but even more in Japan, whose fledgling export industry was struggling to wipe out its reputation for producing low-quality products. The widespread adoption of American ideas of quality improvement has often been cited as an explanation of the so-called Japanese economic miracle, a long

period of economic growth that began after World War II and lasted throughout the 1980s (Bergman 2018, p. 333; Berwick 1989, p. 54). Japanese companies with strict devotion to quality throughout their organization often achieved results that astonished Western visitors:

When a team of Xerox engineers visited Japan in 1979, they discovered that competitors were manufacturing copiers at half of Xerox's production costs and with parts whose freedom from defects was better by a factor of 30. (Abelson 1988)

The word “kaizen,” which means “improvement,” was used in the Japanese quality improvement movement as a general motto, covering all the aspects in which industrial products and processes should be improved. In English, “kaizen” has frequently, but not quite accurately, been translated as “continuous improvement,” abbreviated CI (Singh and Singh 2009). With this terminology, the quality improvement ideas developed in the United States in the 1920s and 1930s were reimported into the United States and other Western countries in the 1980s and 1990s. The main ideas were summarized by a group of British management researchers as follows:

In its simplest form CI can be defined as a *company-wide process of focused and continuous incremental innovation*. Its mainspring is incremental innovation—small step, high frequency, short cycles of change which taken alone have little impact but in cumulative form can make a significant contribution to performance. (Bessant et al. 1994, p. 18)

The concept of continuous improvement also found resonance in the safety professions. However, in spite of its origin in industrial management, it is not in workplace safety but in patient safety that continuous improvement has been most successful. Beginning in the 1980s, professionals working with patient safety in hospitals and clinics around the world have adopted continuous improvement as an overarching idea for their activities (Batalden and Stoltz 1993; Bergman 2018; Berwick 1989, 2008; Berwick et al. 1990). A major reason for its success in healthcare is that continuous improvement fits well into the evaluation culture in modern healthcare, with its use of standardized treatment protocols. Just as in industrial management, continuous improvement in healthcare means that there is no “optimal quality level beyond which further improvement would not be worth the incremental cost of achieving it”; to the contrary, “each instance of improvement is an invitation to consider options for further improvement” (Huckman and Raman 2015, p. 1811). This way of thinking is very much in agreement with Vision Zero and other zero goals.

There is also another interesting similarity between continuous improvement and Vision Zero: They both tend to be closely aligned with the environmental approach to accidents and other untoward events (cf. section “The Zero Subject”). Although doctors and nurses are still personally responsible for the treatments they recommend and administer, the focus in healthcare is shifting away from individual weaknesses to weaknesses in routines, technologies and organizational structures (Kohn et al. 2000). The reason for this is that, perhaps in particular in healthcare,

“defects in quality could only rarely be attributed to a lack of will, skill, or benign intention among the people involved with the processes” (Berwick, 1989, p. 54). The need for an organizational focus in patient safety was explained as follows by two researchers in healthcare management:

M[aintenance] O[f] C[ertification] aims to certify that individuals are proficient in their target responsibilities. This individual certification offers some degree of reassurance to patients who want to know that they are receiving treatment from qualified physicians. Continuous process improvement, however, assumes that quality primarily depends on the process, not simply the individuals who execute it. A central tenet of continuous process improvement is that the problem must be separated from the person. This recognition is important for at least [two] reasons. First, it focuses attention on the process, which is often the root cause of a defect. Second, it makes it safe for workers to highlight issues without concern that they or their colleagues will experience adverse consequences, such as being blamed as the source of the problem. (Huckman and Raman 2015)

In healthcare, continuous improvement has mostly been applied as a principle only for safety. It has usually not been directly aligned with improvements in terms of other goals, such as cost containment or increased productivity. In contrast, the application of continuous improvement in industrial safety is usually part of a more general management strategy, which also includes improvements in other respects than safety. Cost reduction is often a dominant criterion of what constitutes an improvement (Baghel 2005, pp. 765–766). In the literature on continuous improvement, examples are often given that show how productivity, quality, safety and economic output can all be improved at the same time. Potential conflicts among these goals are less often discussed. A prominent exception can be found in a 2002 report from the Nuclear Energy Agency of the OECD on improvements in nuclear plants:

The modern equipment often detects cracks and faults in components and welds that were undetectable by the equipment available when the plant was constructed. If the plant has operated safely and reliably for many years, and there is good evidence that the defect is not “growing”, should the regulator require the defect to be repaired, especially if the repair might degrade other safety features of the plant? Such questions present a real challenge to the regulator when he has to decide how to react to such new information and he must be clear whether he is requiring the licensee to maintain safety or to improve safety. The costs involved can be very great and, in the present financial climate, utilities are likely to mount strong challenges to requirements which they perceive go beyond the original design basis. . .

The lesson for nuclear safety here may very well be: *qui n’avance pas recule!* [Who doesn’t advance retreats.] (Nuclear Energy Agency 2002, pp. 11 and 17)

Similar situations may well occur in other areas and should then be discussed and dealt with in a transparent and responsible manner.

As Low as Reasonably Achievable (ALARA)

The as low as reasonably achievable (ALARA) principle originated in mid-twentieth century radiation protection.

Within 5 years after Röntgen's discovery of X-rays in 1895, researchers working with radioactive material noted that high exposures gave rise to skin burns. It was generally believed that exposures low enough not to produce such acute effects were innocuous, but some physicians warned that radiation might have unknown detrimental effects (Kathren and Ziemer 1980; Oestreich 2014). In the Manhattan project, which developed the first nuclear weapons, the radiologist Robert S. Stone (1895–1966) in the Health Division of the Metallurgical Laboratory in Chicago was assigned to determine “tolerance levels” for radiation exposures of workers. He reported that there was no known safe level for such exposures. Instead of fixed tolerance levels, he proposed that exposures should be kept as low as practically possible. His proposal was accepted (although some wartime exposures were very high, judged by modern standards) (Auxier and Dickson 1983).

After the war, this precautionary no-limit approach was much strengthened by growing awareness that exposure to ionizing radiation increases the risk of leukaemia. The risk appeared to be stochastic. It increases with increasing exposures, but researchers could not identify any “threshold dose,” i.e. any exposure level above zero, below which the risk would be zero. Many scientists believed that the risk of radiation-induced leukaemia was approximately proportionate to the radiation dose (the “linear dose-response model”) (Lewis 1957; Brues 1958; Lindell 1996). In consequence, Robert Stone's approach was adopted by the US National Committee on Radiation Protection (NCRT). In a 1954 statement, they declared that radiation exposures should “be kept at the lowest practical level” (Auxier and Dickson 1983). In 1958, the International Commission on Radiological Protection (ICRP) took a similar standpoint, based on a review of what was then known about the dose-response relationships of leukaemia and other cancers:

The most conservative approach would be to assume that there is no threshold and no recovery, in which case even low accumulated doses would induce leukaemia in some susceptible individuals, and the incidence might be proportional to the accumulated dose. The same situation exists with respect to the induction of bone tumors by bone-seeking radioactive substances. . .

It is emphasized that the maximum permissible doses recommended in this section are *maximum* values; the Commission recommends that all doses be kept as low as practicable, and that any unnecessary exposure be avoided. (ICRP 1959, pp. 4 and 11)

This recommendation has been reconfirmed in a long series of decision by the ICRP. In 1977, it was rephrased as a requirement that “all exposures shall be kept as low as reasonably achievable, economic and social factors being taken into account” (ICRP 1977, p. 3). This principle goes under many names. Apparently, it was first called “as low as practicable” (ALAP), but that name was soon replaced by “as low as reasonably achievable” (ALARA) and “as low as reasonably practicable” (ALARP) (Wilson 2002). Currently, it is known under the following names and abbreviations (Reiman and Norros 2002; Nuclear Energy Agency 2002, p. 14; HSE 2001, p. 8):

As low as practicable (ALAP)

As low as reasonably achievable (ALARA)

As low as reasonably attainable (ALARA)
As low as reasonably practicable (ALARP)
So far as is reasonably practicable (SFAIRP)
Safety as high as reasonably achievable (SAHARA)

ALARA is now recognized worldwide as a principle for radiation protection. It is usually applied to collective doses (i.e. the sum of all individual doses in a plant or an activity), rather than to individual doses. It is therefore often seen as a utilitarian principle. However, it is combined with upper limits for individual exposures, which can be interpreted as based on deontological principles (Hansson 2007a, 2013b). In most countries, this principle is not much used outside of radiation protection. The major exception is Britain, where it has an important role in general worker's health and safety. In this application, it is mostly applied to individual risks (HSE 2001).

In the interpretation of ALARA, it is essential to pay attention to the meaning of the term "reasonable" (alternatively "achievable" or "practicable," in other names of the principle). This term is often used in legal texts such as regulations and rulings, where it has two major functions (Corten 1999). First, it makes regulations adaptable and allows their interpretation to be adjusted to circumstances unforeseen by the lawmaker. In this way, the word "reasonable" can resolve "a contradiction between the essentially static character of legal texts and the dynamic character of the reality to which they apply" (ibid., p. 615). Secondly, references to reasonableness provide legitimacy to a legal order "by presenting an image of a closed, coherent and complete legal system." The notion "masks persistent contradictions regarding the meaning of a rule, behind a formula which leaves open the possibility of divergent interpretations" (ibid., p. 618).

The reasonableness of the ALARA principle (the "R" in the acronym) appears to have both these functions. The first function becomes apparent in the use of "reasonableness" in adjustments to various kinds of economic and practical constraints. It allows the ALARA principle to be "balanced against time, trouble, cost and physical difficulty of its risk reduction measures" (Melchers 2001). The second function shows up when divergences between safety and other considerations are "internalized" within safety management by treating certain potential solutions to safety problems as unreasonable, instead of presenting these divergences as conflicts to be resolved.

According to the original conception of ALARA, it applies to all non-zero risks. For instance, even very low radiation doses should be eliminated if this can be done. However, ALARA has often been reinterpreted so that it only applies to risks above a certain threshold of concern. Hence, the Health and Safety Executive (HSE) in Great Britain has introduced a three-levelled approach to risk, dividing risk exposures into three ranges. In the highest range, risks have to be reduced irrespective of the costs, and there is no need for ALARA considerations. In the lowest range, risks are assumed to be acceptable, which means that there is no need for reducing them, and consequently, ALARA does not apply. It is only in the medium range that ALARA-based activities are said to be applicable. According to the HSE's tentative

limits for the three regions, the “ALARA region” comprises activities with an individual risk of death per year between one in a million and one in one thousand (HSE 2001, pp. 42–46). A similar three-levelled approach has also been proposed for exposures to ionizing radiation (Kathren et al. 1984; Hendee and Edwards 1986).

Such interpretations of ALARA, which disallow its application to risks below a certain level, regardless of how cheaply and easily they can be reduced, make ALARA considerably weaker than full-blown improvement principles such as Vision Zero and continuous improvement. The problems with such a weakening were eloquently expressed already in 1981 by two leading radiation protection experts, Bo Lindell (1922–2016) and Dan J. Beninson (1931–1994):

[I]n each situation, there is a level of dose below which it would not be reasonable to go because the cost of further dose reduction would not be justified by the additional eliminated detriment. That level of dose, however, is not a *de minimis* level below which there is no need of concern, nor can it be determined once and for all for general application. It is the outcome of an optimization assessment which involves marginal cost-benefit considerations. . . It is not reasonable to pay more than a certain amount of money per unit of collective dose reduction, but if dose reduction can be achieved at a lesser cost even at very low individual doses, the reduction is, by definition, reasonable. (Lindell and Beninson 1981, p. 684)

Best Available Technology (BAT)

Probably the first legal requirement to employ the best practicable technology to solve environmental problems can be found in the British Alkali Act of 1874. After specifying in some detail how emissions of hydrogen chloride should be reduced, the Act continued:

In addition to the condensation of muriatic acid gas [hydrogen chloride] as aforesaid, the owner of every alkali work shall use the best practicable means of preventing the discharge into the atmosphere of all other noxious gases arising from such work, or of rendering such gases harmless when discharged. (Anon 1874, p. 168)

The term “best practicable means” was interpreted as referring to both economic limitations and technological feasibility (Holder and Lee 2007, p. 331).

When new and more ambitious environmental regulations were introduced in the 1960s and 1970s, it was soon discovered that statutes requiring specific technological solutions have two important problems: They are *unadaptable*, since they do not allow industry to achieve the same effect with different technical means, and they are also *slow-moving* and tend to lag behind when new and better technological solutions become available. Regulations specifying maximal allowed emissions are more adaptable, since they leave it to industry to decide how to achieve the required emission standard. However, they are just as slow-moving as statutes requiring a specific technology (Sunstein 1991, pp. 627–628n; Ranken 1982, p. 162). Legislation based on the best available technology (BAT) was introduced as a way

to solve both these problems and make legislation both adaptable and sufficiently fast-moving. The basic idea was to require use of the best available emission-reducing technology. Such a rule is technology neutral; if there are alternative ways to reach the best result, then each company can make its own choice among these alternatives. Furthermore, BAT statutes can stimulate innovations in environmental technology. If a new technological solution surpasses those previously available, then it becomes the new BAT standard to which industry must adjust.

A large number of synonyms and near-synonyms of “best available technology” have been used in different legislations, including the following (Merkouris 2012; Vandenbergh 1996; Ranken 1982):

- Best available control technology (BACT)
- Best available techniques (BAT)
- Best available technology not entailing excessive costs (BATNEEC)
- Best environmental practice (BEP)
- Best practicable control technology (BPT)
- Best practicable environmental option (BPEO)
- Best practicable means (BPM)
- Lowest achievable emissions rate (LAER)
- Maximum achievable control technology (MACT)
- Reasonably achievable control technology (RACT)

Several of the above terms, perhaps in particular LAER, can also be interpreted as variants of the ALARA principle. In some legislations, more than one of these concepts are used, often with different specifications. For instance, in American legislation, “best practicable control technology” (BPT) has been used for lower demands on emissions control than the stricter “best available technology” (BAT).

In the United States, BAT strategies were introduced in most environmental legislations in the 1970s and 1980s and became “a defining characteristic of the regulation of the air, water, and workplace conditions” (Sunstein 1991, pp. 627–628). Legislation based on the BAT concept has also been introduced in most European countries and in legislation on the European level. The European Directive on Industrial Emissions does not allow large industrial installations to operate without a permit that imposes emission standards based on best available techniques (Merkouris 2012; Schoenberger 2011). The BAT concept is also employed in several international treaties, such as the 1992 Convention on the Protection of the Marine Environment of the Baltic Sea Area (the Helsinki Convention) and the Convention for the Protection of the Marine Environment of the North-East Atlantic (the OSPAR Convention) from the same year (Merkouris 2012).

BAT requirements are widely used in emissions control. However, there is one major category of emissions for which they are not much used: Legislation on the limitation and reduction of greenhouse gas emissions has in most cases been based on other regulatory principles, including standards based on current technologies and tradable emission permits. Perhaps surprisingly, the BAT concept does not either seem to have been used systematically in safety legislation. For instance, type

approval and similar procedures for motor vehicles, aircraft, electrical appliances, etc. are based on well-defined technical standards, and there does not seem to be a movement towards replacing such standards by reference to the best available technology. Proposals have sometimes been made to apply the BAT concept to other areas than emissions control, but with relatively little success (Helman and Parchomovsky 2011). One of the few cases in which BAT principles have been applied to safety is the 1967 guidelines for the safety standards to be developed by the US National Highway Traffic Safety Administration. Such standards were to be “stated in terms of performance rather than design specifying the required minimum level of performance but not the manner in which it is to be achieved” (Blomquist 1988, p. 12).

Just as the stringency of ALARA principles depends on how the R (“reasonable”) is interpreted, the stringency of BAT principles hinges on the interpretation of the A (“available”). Originally, BAT regulations did not require a cost-benefit analysis. This has often been an advantage from the viewpoint of safety. For instance, this made it possible for American regulators to ensure that offshore technologies “truly implement the best available technology as opposed to technology that is only economically convenient” (Bush 2012, p. 564). However, in some BAT regulations, “available” is interpreted as economically feasible. This has led to the linguistically somewhat awkward situation that there are technologies on the market that are better than the “best available technology” (but too expensive). Such technologies have been called “beyond BAT” (Schoenberger 2011).

On the other hand, there are cases in which not even the best technology that is at all available (at any price) is good enough to protect the environment. In some such cases, regulatory agencies have been authorized to impose requirements stricter than the BATs in order to achieve sufficient protection of the environment (Vandenbergh 1996, pp. 837–838 and 841). Some authors claim that BAT regulations are inadequate since they put focus on what can currently be done rather than on what is most important to do, thereby distracting “from the central issue of determining the appropriate degree and nature of regulatory protection” (Sunstein 1991, p. 629; cf. Ackerman and Stewart 1988, pp. 189–190).

Summary

The improvement principles that we have studied in this and the previous section have much in common. They all carry the same basic message, namely, that as long as improvement in safety is possible, it should be pursued. They therefore serve as antidotes to fatalism and complacency.

None of this disallows economic and other competing considerations from having a role in determining the pace and means of implementation. It would be futile to prescribe that all improvements should be implemented immediately, regardless of costs. But there is an important difference between postponing a safety improvement and dismissing it altogether. When compromises with other social objectives are necessary, the improvement principles induce us to see these compromises as

temporary, unsatisfactory concessions. This is a clear signal that currently unrealistic safety improvements should be pursued if and when they become realistic and that innovations that put them within reach are most welcome.

However, several of the improvement principles have been subject to reinterpretations that dismiss instead of postpone currently unrealistic safety improvements. For ALARA, this has taken the form of interpreting the R (“reasonable”) as excluding the reduction of comparatively small risks, even if such reductions can be done with very small effort and sacrifice. For BAT, the A (“available”) has been reinterpreted so that costly but affordable reductions in emissions are not required. It is one of the advantages of Vision Zero and other zero goals that they do not easily lend themselves to such debilitating reinterpretations.

Aspiration Principles

The improvement principles form part of a larger group of safety principles, namely, those that tell us what levels of safety or risk reduction we should aim at or aspire to. This larger group can be called the *aspiration principles* (Hansson 2019). As shown in Fig. 1, we can distinguish between three major types of aspiration principles, in addition to the improvement principles.

Acceptance principles draw a line between acceptable and unacceptable risks. That limit depends on the risks alone, without taking the benefits that come with the risks into account. We will consider three types of acceptance principles: *risk limits*, *exposure limits* and *equipment and process regulations*.

Weighing principles require that we weigh safety against other objectives, such as productivity and economic gains, and strike a balance between them. Whereas acceptance principles usually have an affinity with deontological (duty-based) moral thinking, weighing principles have much in common with consequentialist ethics. We will discuss three types of weighing principles, namely, *cost-benefit analysis*, *individual cost-benefit analysis* and *cost-effectiveness analysis*.

Finally, we will discuss *hypothetical retrospection*, a safety principle requiring that our decisions will be defensible also in the future.

Risk Limits

In the early days of risk analysis, some risk analysts maintained that all dangers falling below a certain *risk limit* are acceptable. That limit was usually expressed as a probability of death, often a “cut-off level of 10^{-6} individual lifetime risk [of death]” (Fiksel 1985, pp. 257–258). That idea has now largely been replaced by more sophisticated approaches that weigh risks against the benefits they are accompanied by. However, the idea of a risk limit has repeatedly been revived, usually under the auspices of a “de minimis” position in risk regulation, according to which there is a probability threshold, below which a risk is always acceptable even if it comes without any advantages.

It does not take much intellectual effort to see that this is an untenable approach (Pearce et al. 1981; Bicevskis 1982; Otway and von Winterfeldt 1982; Hansson 2013a, pp. 97–98). To begin with, since the “de minimis” principle is applied to each risk individually, it does not protect us against large cumulative effects of a large number of risks, each of which falls below the limit. For instance, in modern societies, we are exposed to a large number of chemical substances. If each of them were to give rise to a “de minimis” risk, the combination of them all could nevertheless be far above the risk limit.

More fundamentally, even risks with very low probabilities are clearly unjustified if they bring nothing good with them. For a simple example, suppose that someone constructs a bomb connected to a random generator, such that the probability is 10^{-9} that the bomb will detonate. The bomb has been covertly mounted in a place where it will kill exactly one (unsuspecting) person if it explodes. The risk associated with such a device is “de minimis” according to the usual criteria. However, it is clearly unacceptable, for the simple reason that it imposes a frivolous, completely unjustified risk on a person who did not ask for it. On the other hand, we routinely take risks of 10^{-9} or higher in order to gain some advantage. Travelling an hour by car is one example of this (International Transport Forum 2018, p. 21). We can conclude from examples like this that the acceptability of a risk imposition cannot be determined based only on the size of the risk. Other factors, such as the associated benefits, have to be taken into account, even if the risk is very small. Reliance only on the size of the risk has been called the “sheer size fallacy” in risk analysis (Hansson 2004a, pp. 353–354).

Exposure Limits

Our second group of acceptance principles is *exposure limits*, numerical upper bounds on allowable exposures to chemical substances and to physical hazards such as noise and radiation. Exposure limits can be based on the same types of considerations as risk limits, but they are much easier to implement, since exposures can usually be measured with well-established chemical and physical methods.

The first limits for occupational exposures were proposed by individual researchers in the 1880s. In the 1920s and 1930s, several lists were published in both Europe and the United States, and in 1930, the USSR Ministry of Labor issued what was probably the first official list (Cook 1987, pp. 9–10). In 1946, the American Conference of Governmental Industrial Hygienists published the first edition of their list. This list, which is revised yearly, has long been a standard reference for official lists all over the world. Since the 1970s, most industrialized countries have their own lists of occupational exposure limits. Exposure limits for ambient air were introduced in the same period (Greenbaum 2003). In food safety, exposure limits were introduced in the 1960s under the name of “acceptable daily intake” (Lu 1988).

The first proposed exposure limit for ionizing radiation was published in 1902. It aimed to protect against the acute effects, but it was based on a rather primitive and

unreliable method of measurement. In the 1920s and 1930s, improved methods for dose measurements were developed and put to use in the implementation of more precise exposure limits. In the 1950s, exposure limits were adjusted to take the long-term carcinogenic effects of ionizing radiation into account (Parker 1980, pp. 970–971; Broadbent and Hubbard 1992).

Ideally, one might hope that exposure limits should guarantee safety, in the sense that exposures below the limits impose no risk. Unfortunately, that is often not the case, for three major reasons. First, many standards, in particular those for occupational exposures, result from compromises with economic considerations. This has often led to exposure limits at levels that are known to be associated with occupational disease (Hansson 1997b, 1998a, b; Johanson and Tinnerberg 2019). These risks can be considerable for the average worker, but they are even greater for workers who are particularly sensitive, for instance, due to pregnancy or prior disease (Johansson et al. 2016; Hansson and Schenk 2016). Secondly, long-term effects of chemical exposures are difficult to determine, and some exposure limits that were believed to be safe have later been shown to be unsafe due to previously unknown effects of the substance. One example of this is the drastic reduction of the exposure limit for vinyl chloride from 500 ppm (parts per million) to 1 ppm when the carcinogenicity of this substance was discovered in 1974 (Soffritti et al. 2013). Thirdly, for many carcinogenic substances, it is impossible to determine a risk-free exposure level above zero. The best estimate seems to be that the risk of cancer is proportional to the exposure, which means that every non-zero exposure limit is associated with an implicit level of accepted risk. For all these reasons, exposure limits should not be considered as safe limits. Gains in safety can be expected if exposures are reduced as far below current exposure limits as possible.

Exposure limits are constructed to have a very wide application. Occupational exposure limits for chemical substances apply to all workplaces where the substances are used. Similarly, air quality standards for ambient air apply to outdoor air everywhere in the jurisdiction. This general applicability is unproblematic for health protection if the exposure limit represents a level below which there are no adverse effects on the exposed population. However, if the limit represents a compromise between health protection and economic considerations, then the general applicability tends to lead to exposure limits that are unnecessarily high in many of the places where they apply. This is because economic considerations at the places where exposure reductions are expected to be most costly tend to dominate the standard-setting process. A classic example of this is the exposure limit of 1 ppm for the carcinogenic substance benzene that was adopted by the US Occupational Safety and Health Administration in 1987. Values lower than this were considered infeasible due to excessive compliance costs in the petrochemical, coke and coal industries. However, only 2.2% of the workers exposed to benzene worked in these industries (Rappaport 1993, p. 686). The remaining 97.8% of American workers exposed to benzene, about 230,000 workers, had a weaker protection against benzene than what would have been economically feasible in their own branches of industry.

It is not unreasonable to ask: If workers in a particular industry have to be exposed to high levels of a hazardous substance, is that really a reason to accept equally high

levels in other industries where it would be comparatively easy and inexpensive to comply with a considerably lower exposure limit? An alternative approach in this situation would be to adopt a lower general exposure limit that is realistic on most workplaces, in combination with regularly reviewed, higher, exception values for branches of industry that are not yet capable of complying with the general value (Hansson 1998a, pp. 106–109).

Process and Equipment Regulations

Our third group of acceptance principles is *process and equipment regulations*. Regulations requiring machines to be equipped with certain safety features have a long tradition, in particular in occupational safety. In Britain, already the pioneering Factories Act of 1844 contained stipulations on both equipment and work processes. All mill gearing, as well as certain other moving parts of machines, had to be securely fenced. Furthermore, it was prohibited to use children or young workers to clean the mill gearing while it was in motion (Hutchins and Harrison 1911, pp. 85–87; Tapping 1855, pp. 43–47). This was the beginning of increasingly strict regulations on equipment and processes, which have contributed much to the reduction of many types of workplace injuries. Hand injuries from mechanical power presses are among the best known examples.

Such regulations have been equally important in road vehicle safety. Since its beginnings in the late nineteenth century, the legislation on motor vehicles has developed gradually from an almost exclusive focus on driver behavior to increasingly strict requirements on vehicle construction. For instance, the first British legislation on motor cars was the Locomotives on Highways Act of 1896, in which motor cars were called “light locomotives.” That legislation was focused on the behavior of drivers, who were required to have a license, which could be suspended in case of misconduct. A general speed limit of 14 mph (23 km/h) applied to all “light locomotives.” In the detailed regulations based on the Road Traffic Act of 1930, much more emphasis was put on the construction of vehicles, which were, for instance, required to have unimpaired view ahead, safety glass in windcreens and rear-view mirrors. This was followed in 1937 by requirements for windscreen wipers and speed indicators (Tripp 1938). In the 1960s, important further steps were taken in many countries towards making vehicles safer (Furness 1978). The United States had an important role in this development. The National Traffic and Motor Vehicle Safety Act of 1966 introduced a new way of thinking about traffic safety. A federal agency, the National Highway Traffic Safety Administration, which is still in operation, was created with the explicit task to make manufacturers produce vehicles with reduced risk of crashes and improved protection of the occupants of the vehicle in case of a collision (Mashaw and Harfst 1987; Blomquist 1988). The general approach taken by the administration was to reduce traffic casualties as much as possible. This was noted by economist Glenn Blomquist in a book criticizing their approach:

Each year the NHTSA [National Highway Traffic Safety Administration] prepares a report on its activities under the Vehicle Safety Act. Each year changes in the number of traffic fatalities and in the fatality rate (per vehicle miles) are described. Some years the fatalities and rates are up and some years the fatalities and rates are down compared to previous years. Every year, however, the implication is the same: the traffic safety problem deserves more attention than ever before. Travel risks are not zero despite effective policy is the contention. The reasoning seems to be if fatalities and rates are up then more aggressive policy is needed to bring them down, and if fatalities and rates are down, then more aggressive policy is needed to reduce them further. (Blomquist 1988, pp. 115–116)

According to Blomquist, this showed that the NHTSA entertained a “risk-free goal.” In his view, such a goal is “unwise and futile” since “no agency will ever have sufficient power or resources to completely control individual behavior” (*ibid.*, p. 115). There is of course another side to this; with a more positive view on zero goals, the NHTSA can instead be described as a forerunner of a modern, more progressive approach to technology improvement.

Today, the development of safety standards for motor vehicles is largely driven through international cooperation, in which the World Forum for Harmonization of Vehicle Regulations has a central role. In this and other areas, technological safety regulations tend to be quite specific on what is required, and best available technology (BAT) clauses are seldom if ever used. However, in areas such as motor vehicle safety where regulators actively follow the technology development in detail, technological improvements can still be introduced by timely amendments of regulations.

Cost-Benefit Analysis

We will now turn to the next main category of aspiration principles, namely, weighing principles. These are principles demanding the weighing of safety objectives against various other objectives with which they may run into conflict, such as ease of work, product quality, environmental protection, productivity, cost containment and economic gains. Most of the discussion has focused on conflicts between safety and economic limitations, but in practice, safety concerns can also clash with various non-economic constraints and objectives. The dominant weighing principle is cost-benefit analysis (CBA). It is a powerful economic tool, but it is based on simplifying assumptions that are far from unproblematic.

The basic idea of cost-benefit analysis is quite simple: In order to compare the advantages and disadvantages of decision alternatives, they are all assigned a monetary value. Suppose that a proposed new road project costs 25 million euros. Furthermore, it is expected to lead to a total reduction in traffic time for all its users of 6,000,000 hours and the loss of four unique local species of hoverflies. We assign the value of 5 euros to each gained hour and the value of 1 million euros for each hoverfly species. If these are the only factors to be taken into account, then the total value of the project is as follows:

$$5 \times 6,000,000 - 4 \times 1,000,000 - 25,000,000 = 1,000,000 \text{ euros.}$$

Since the total value is positive, the analysis recommends that the road be built. A major problem in this example is of course how to determine the economic values of travel time and lost species. Proponents of cost-benefit analysis emphasize that since we do not have unlimited amounts of money, there is no way to avoid weighing non-monetary values against monetary costs. For instance, we take measures to save a hoverfly species if doing so does not cost much, but we will not do it at any price. According to the proponents of cost-benefit analysis, the major difference is that with this method, we make these decisions transparently, basing our decisions on known prices, rather than unarticulated intuitions. If we use the same monetary values in different decisions, then we can also achieve increased consistency in our decision-making processes.

When cost-benefit analysis is applied to safety decisions, uncertain outcomes will have to be included in the analysis. This is usually done by assigning to each such outcome the best available estimate of its expectation value (probability-weighted value). For instance, suppose that 200 deep-sea divers perform an operation in which the risk of death is 0.001 for each individual. Then, the expected number of fatalities from this operation is $0.001 \times 200 = 0.2$. If we apply a “value of life” of 3 million euros, then the monetary cost assigned to this series of dives is 0.6 million euros.

Cost-benefit analysis involves a rather radical simplification of multidimensional real-life problems in order to make them accessible to a transparent and easily manageable one-dimensional analysis. Unsurprisingly, this gives rise to a host of philosophical and interpretational issues (Hansson 2007c). Here, we will focus on four problems that are highly relevant for safety applications, namely, incommensurability, incompleteness, collectivism and complacency.

By incommensurability between two values is meant that they are so different in nature that no translation between them is possible. Probably the most common criticism of cost-benefit analysis is that it violates the incommensurability between human life and money. The assignment of a monetary value to human lives is said to violate the sanctity of life (Anderson 1988; Sagoff 1988; Hampshire 1972, p. 9). One reason why this criticism is so widespread may be that cost-benefit analysts have failed to explain the difference between the calculation values used in their analyses and prices on a market. The assignment of a sum of money to the loss of a human life does not imply that someone can buy another person, or the right to kill her, at that price. A more serious problem may be the arbitrariness of the values used in cost-benefit analyses. Not only do we lack a well-founded answer to what calculation value should be used for the loss of a human life. We also lack definite answers to questions such as how many cases of juvenile diabetes correspond to one death or what amount of human suffering or death corresponds to the extinction of an antelope species. Methods have been developed to determine monetary values for these and other seemingly non-monetary assets, but these methods are all fraught with uncertainty, and none of them has a reasonably sound philosophical foundation (Heinzerling 2000, 2002; Hausman 2012).

By incompleteness is meant in this context that factors that could legitimately have an influence on a decision are left out of the analysis. Even quite extensive cost-benefit analyses of societal projects tend to leave out decision effects that are difficult to express in quantitative terms. This applies, for instance, to risks of cultural impoverishment, social isolation and increased tensions between social strata. Such issues may nevertheless be important considerations for decision-makers. Unfortunately, there is often a trade-off between attempted solutions to the incompleteness problem and incommensurability problems. In order to solve incompleteness, we would have to assign monetary value to additional potential effects, such as social incohesion, which are extremely difficult to monetize. But by doing so, we would aggravate the problem of incommensurability.

The collectivism of standard cost-benefit analysis is a consequence of its aggregation of all effects to a single number, irrespectively of whom they accrue to. In our above example of a road project, the reduction in travel time was judged by the total sum for all travellers, 6,000,000 h. The distribution of these gains has no influence on the analysis. This net gain could, for instance, arise as a result one million long-distance travellers gaining 7 h each, whereas each of five thousand local travellers has to spend 200 h more travelling. This would be very different from a situation where only one hundred thousand travellers were affected, and they all gained 60 h each (Nordström et al. 2019). Standard cost-benefit analysis makes no difference between these two situations, since the total net effect on travel time is the same. Even worse, cost-benefit analysis treats serious risks such as death risks in the same way. Distributional issues are simply not part of its standard considerations. This is of course particularly problematic if the benefits and the disadvantages of a project are received by different groups of people.

The complacency induced by cost-benefit analysis consists in its tendency to foster acceptance of those evils that cannot currently be rectified with a positive cost-benefit analysis. For an example, suppose that a country has a large number of unguarded railroad crossings in thinly populated areas. Each year, several fatalities are caused by collisions between trains and vehicles or pedestrians passing one of these crossings. The number of fatalities can be drastically reduced by installing traffic lights and half-barrier gates, operated by the rail traffic control system. However, this would be much too expensive, due to the large number of crossings in places with few road users. The message of a cost-benefit analysis in such a situation would be that the life-saving traffic control system is simply not optimal and cannot be defended. In contrast, the message emerging from Vision Zero or other improvement principles would be that the life-saving system is indeed desirable but cannot be implemented at present, due to other, even more pressing priorities. The latter message has the obvious comparative advantage of being more conducive to cost-reducing innovations and more promotive of continued social activities in the issue.

But as already mentioned, there are other ways to balance advantages against disadvantages. In the next two subsections, we will briefly consider two alternatives to standard cost-benefit analysis.

Individual Cost-Benefit Analysis

In *individual cost-benefit analysis*, costs and benefits affecting different individuals are not added up. Instead, a separate cost-benefit analysis is made for each individual or, in practice, for each type of concerned individual (Hansson 2004b). For instance, in a road project, separate cost-benefit analyses can be made for categories such as local inhabitants, people driving a private car to and from work on the road and people travelling daily on it in buses. The outcomes of these different cost-benefit analyses will typically differ, and they may even point in different directions concerning the value of the project. This should not be seen as a disadvantage. A necessary first step towards solving conflicts of interest is to recognize them.

Cost-Effectiveness Analysis

The other, probably more important, alternative to standard cost-benefit analysis is *cost-effectiveness analysis* (CEA), which compares costs and benefits by calculating cost-effect ratios. For instance, if the desired effect of a technological innovation in motor cars is to reduce the number of fatalities, then the outcome can be reported as the expected cost per life saved by introducing the innovation in question.

Cost-effectiveness is mostly applied to healthcare interventions, where the most commonly used ratios are (i) cost per life-year gained and (ii) cost per quality-adjusted life-year gained. (The number of saved quality-adjusted life-years is the product of the number of saved life-years with a factor that is 1 if these are years lived in good health but smaller if they are years lived with a severe medical condition.) For instance, a French study investigated the costs and effects of smoking cessation counselling and treatment. It showed an average expected cost of less than 4000 euros per life-year gained (Cadier et al. 2016). Other studies of smoking cessation give similar results. This is an unusually low cost for a life-saving medical intervention, and smoking cessation is therefore an unusually cost-effective medical intervention.

Cost-effectiveness studies are comparatively uncommon outside of the healthcare sector, but there are plenty of examples showing their usefulness. For instance, in studies intended to guide energy saving in buildings, it is highly useful to calculate the cost per kWh energy saved with different energy efficiency measures (Tuominen et al. 2015). Houseowners and other decision-makers can then obtain maximal energy savings for their money by giving priority to the most cost-efficient measures. The cost-efficiency approach appears to be much more appropriate in this case than a cost-benefit analysis, which would divide the measures into two classes, those approved and those disapproved. Safety measures can also be evaluated in this way. For instance, one study showed that engineering control programmes to reduce silica exposure on workplaces are highly cost-effective; some such measures had a cost of only about USD 110 per quality-adjusted life-year (Lahiri et al. 2005; cf. Tengs et al. 1995).

Cost-effectiveness analysis is eminently useful when a proposed safety measure has to be evaluated in terms of its effects on safety and its cost, and no additional factors need to be taken into account. The application of this method is much less clear-cut if there are also other effects, say effects on the environment, that have to be taken into account. It can also be difficult to apply if the safety measure is an integrated part of some larger project. For instance, if a new road is built to replace an old unsafe road with too little capacity, then it is usually not possible to divide up the project costs between costs for increased safety and costs for increased capacity. But in the cases with well-defined costs for safety, cost-effectiveness analysis has distinct advantages over cost-benefit analysis and should probably be used more often.

Hypothetical Retrospection

Safety management is largely a matter of giving sufficient weight to untoward events that might happen in the future. A basic type of reasoning to that effect is the “foresight argument” (Hansson 2007b, p. 147). It urges us to take into account the possible effects of what we do now on what can happen later. The argument has both a deterministic and an indeterministic variant. As an example of the deterministic variant, some of the consequences of drinking excessively tonight can, for practical purposes, be regarded as foreseeable. As an example of the indeterministic variant, driving drunk substantially increases the risk of causing an accident, but of course, there is also a considerable chance that nothing serious will happen. Nevertheless, the increased risk is reason enough not to drink and drive.

The indeterministic variant of the foresight argument is highly useful for thinking about safety. It requires that we think through the various ways in which the future can develop and pay special attention to those “branches” of future development in which things go seriously wrong. It can therefore be described as the very antithesis of wishful thinking. Its purpose is to ensure, as far as possible, that whatever happens in the future, it will not give us reason to say that what we do now was wrong. To achieve this, we can systematically consider what we plan to do now from alternative future perspectives. This way of thinking is called *hypothetical retrospection*.

This may seem difficult and perhaps overly abstract, but it is in fact a way of thinking that we teach our children when trying to help them become responsible and thoughtful persons. “Do not leave all that homework to tomorrow! You know very well how you will feel tomorrow if you do so.” “Save some of the ice-cream for tomorrow. You know that you will regret if you don’t do it.” And of course, as grown-ups, we sometimes wish that we had been more proficient at “thinking ahead” about various subject matter. To apply hypothetical retrospection in safety management means to methodically develop and apply that way of thinking in one’s area of professional responsibility.

The following example exemplifies what this can mean in practice:

A factory owner has decided to install an expensive fire alarm system in a building that is used only temporarily. When the building is taken out of use, the fire alarm has yet never

been activated. The owner may nevertheless consider the decision to install it to have been right, since at the time of the decision other possible developments (branches) had to be considered in which the alarm would have been life-saving. This argument can be used not only in actual retrospection but also, in essentially the same way, in hypothetical retrospection before the decision. Similarly, suppose that there is a fire in the building. The owner may then regret that she did not install a much more expensive but highly efficient sprinkler system. In spite of her regret, she may consider the decision to have been correct since when she made it, she had to consider the alternative, much more probable development in which there was no fire but the cost of the sprinklers had made other investments impossible. (Hansson 2013a, pp. 68–69; cf. Hansson 2007b, pp. 148–149)

For most practical purposes, the application of hypothetical retrospection in safety management consists in following a simple rule of thumb: “Make a decision that you can defend also if an accident happens.” The application of this principle will typically support strivings for risk reduction, and it is therefore concordant with zero goals and other improvement principles.

Summary

In this section, we have studied various aspiration principles, other than the improvement principles that were the topics of the two previous sections. Some of the principles discussed in this section tend to run into conflict with the improvement principles, since they support acceptance of conditions in which safety can still be improved. This applies in particular to cost-benefit analysis, risk limits and exposure limits. These principles share a major problem: they tend to support the presumption that the compromises that have been made (perhaps for good reasons) between safety and other social goals represent a satisfactory state of affairs, thus downplaying the need for future enhancements that go beyond them.

On the other hand, two of the aspiration principles that we have studied in this section, namely, cost-effectiveness analysis and hypothetical retrospection, are easily compatible with the improvement principles. Cost-effectiveness analysis, in particular, can serve as a priority-setting tool to support the application of Vision Zero, continuous improvement and other improvement principles. In a situation with limited economic resources, cost-effectiveness analysis can act as a pathfinder, helping to identify the largest improvements in safety that can be achieved as the next step.

Error Tolerance Principles

Two of the most important insights in safety engineering and safety management are that *things go wrong* and that *humans make mistakes* however much we try to avoid it. Therefore, it is not sufficient to reduce the risk of failures as much as we can. We also have to ensure that the consequences of failures are as small as possible. This is

not a new insight. More than 500 years ago, Leonardo da Vinci (1452–1519) wrote as follows in one of his notebooks:

In constructing wings one should make one cord to bear the strain and a looser one in the same position so that if the one breaks under the strain the other is in position to serve the same function. (Hart 1962, p. 321)

Some of the most important safety principles recommend that equipment, procedures and organizations be so constructed that failures have as small negative consequences as possible. We can call them error tolerance principles. In this section, we will have a close look at six such principles: fail-safety, inherent safety, the substitution principle, safety factors, multiple safety barriers and redundancy.

Fail-Safety

An equipment or procedure is *fail-safe* if it can “fail safely,” which means that the system is kept safe in the case of a failure. Fail-safety can refer to two types of failure: device failure and human failure. The requirements of a fail-safe system have been usefully summarized as follows:

The basic philosophy of fail-safe structures is based on:

- (i) the acceptance that failures will occur for one reason or another despite all precautions taken against them.
- (ii) an adequate system of inspection so that the failures may be detected and repaired in good time.
- (iii) an adequate reserve of strength in the damaged structure so that, during the period between inspections in which the damage lies undetected, ultimate failure of the structure as a whole is remote. (Harpur 1958)

The safety valve is a classic example of a design that makes a system fail-safe. Safety valves are mounted on pressure vessels in order to prevent explosions. (Other means to achieve the same effect are rupture disks, also called burst diaphragms, which act as one-time safety valves, and leak-before-burst design, by which is meant that a crack will give rise to pressure-releasing leakage rather than an explosion.) The origin of the safety valve is not known with certainty, but it is usually credited to the French physicist and inventor Denis Papin (1647–1713), in whose book from 1681 on pressure cookers it was first described (Papin 1681, pp. 3–4; Stuart 1829, p. 84; Le Van 1892, pp. 10–11). In the eighteenth century, safety valves became a standard feature of steam engines. However, enginemen soon found that they could be used to control the machine. Safety valves were frequently tied down or loaded with heavy objects in order to increase the working pressure. These work practices resulted in serious accidents (Hills 1989, p. 129). To prevent such calamities, engine makers provided steam engines with two safety valves. One of them could be operated by the enginemen, whereas the other was inaccessible to them. It could, for instance, be contained in locked, perforated box. This was common practice at the beginning of

the nineteenth century (Partington 1822, pp. 80, 88, 90, 98, 100, 106, 107–108, 109, 114, 115, 116, 122 and Appendix, p. 76). In 1830, an American railway company applied it as a safety rule for their locomotives:

There must be two safety valves, one of which must be completely out of the reach or control of the engine man. (Thomas 1830, p. 373)

Notably, the requirement of a tamperproof safety valve is an early example of a construction tailored to protect not only against machine failures but also against mistakes by the operators.

Another classic example of a fail-safe construction is the so-called dead man's handle (dead man's switch), a control device that has to be pressed continuously in order to keep a machine going or a vehicle moving. The term "dead man's handle" was used already in an American engineering magazine in 1902. The author emphasized that the motorman could only drive the train if he held the handle at all times "and should he drop dead or become disabled, the train will stop of itself, and will not run wild" (Anon. 1902). (However, the handle might not be released if the driver fell over it. Therefore, more advanced vigilance systems are used in modern trains.) Today, similar mechanisms can be found on lawnmowers and on handheld machines such as drills and saws.

A similar mechanism, triggered by device failure rather than human failure, was introduced in the early 1850s by the American inventor Elisha Otis (1811–1861) in his so-called safety elevator. The elevator car was equipped with brakes that automatically gripped the vertical guide rails if the tension of the cord was released, for instance, in the event of a cord break. This invention made elevators safe enough for general use, and it was one of the technical preconditions for the building of skyscrapers that began in the 1880s.

A fail-safe system should go to a safe state in the event of failure. However, technical systems differ in what that safe state is. Trains, lawnmowers, elevators and handheld electric drills can be made safe (or as safe as possible) by being stopped. In all these cases, fail-safety is achieved with a negative feedback that stops movement in the system if a failure occurs. The same applies to a nuclear reactor, in which dangerous conditions should lead to an automatic shutdown. In all these cases, the system is fail-safe if it is *fail-passive* (*fail-silent*) (Hammer 1980, p. 115). However, there are also technical systems in which safety requires normal operations to continue as long as possible even in the event of failure. This applies, for instance, to airplanes. In such cases, a fail-safe system should be *fail-operational* (*fail-active*). This is achieved if the device is sturdy enough to fulfil its function for a sufficient time after it is damaged. This is called *fault tolerance* (*damage tolerance*) and is often achieved with the help of *safety factors* or with *redundancy*, i.e. the duplication of vital components or functions.

Several alternative terms are used for fail-safety when it is primarily aimed to protect against human failures. A system is said to be *foolproof* (*idiot-proof*) if nothing dangerous happens when it is used incorrectly, tampered with or used in unintended ways. Design making a system foolproof is often called *defensive design*.

The Japanese term *poka-yoke* means mistake-proof. It is often used about constructions that prevent human mistakes from leading to product defects, rather than to safety problems.

In 1974, W.C. Clark proposed a distinction between the two terms safe-fail and fail-safe (Jones et al. 1975, p. 1n.). According to this proposal, “fail-safe policy strives to assure that nothing will go wrong,” whereas “safe-fail policy acknowledges that failure is inevitable and seeks systems that can easily survive failure when it comes” (Jones et al. 1975, p. 2). However, these definitions do not correspond to common linguistic practice. The term “safe-fail” is seldom used, and what Jones and co-workers called by that name is usually called “fail-safe.” What Jones and co-workers called “fail-safe” is designated by other terms, such as “inherent safety.”

Inherent Safety

By inherent safety is meant that untoward events are eliminated or made impossible. This contrasts with fail-safety, which reduces the negative effects of untoward events, rather than preventing them from happening. For a simple example, consider a process in which inflammable materials are used. If we replace them by non-inflammable materials, then we have achieved inherent safety. If we still have them but have reduced the consequences of a fire, for instance, by keeping them in containers at safe distance from all buildings, then we have achieved fail-safety.

This distinction has a long tradition. Around 1950, it became common to use the term “primary prevention” for measures against a disease that have the effect of “keeping it from occurring” and “secondary prevention” for “halting the progression of disease after early diagnosis” (Sabin 1952, p. 1270). These terms were soon adopted in accident prevention. In an article in an international road safety journal in 1961, the influential Norwegian civil servant Karl Evang wrote that the concepts of primary prevention (prevention of occurrence) and secondary prevention (prevention of progress) “have now been generally accepted in the field of preventive medicine.” He proposed that they should also be used in the area of traffic safety (Evang 1961, p. 42n).

“Primary prevention” is essentially a synonym of “inherent safety” and “secondary prevention” a synonym of “fail-safety.” The phrase “inherent safety” has been used at least since the 1920s (Bouton 1924), but it acquired its modern sense in the discussions that followed after the disastrous explosion in a chemical plant in Flixborough in June 1974, which caused the death of 28 persons and seriously injured 36. Trevor Kletz (1922–2013), a chemist working for one of the large chemical companies, showed that the accident would not have reached its catastrophic proportions if simple measures had been taken to reduce the hazards. Perhaps most notably, large quantities of inflammable chemicals had been stored close to occupied buildings. Based on these tragic experiences, Kletz proposed that whenever possible, the chemical industry should eliminate hazards rather than just try to manage them. He originally used the term “intrinsic safety” for this concept but soon replaced it by “inherent safety” (Kletz 1978). Four major types of measures are

included in the concept of inherent safety that he and other safety professionals in the chemical industry have developed (Khan and Abbasi 1998; Bollinger et al. 1996):

Minimize (intensify): use smaller quantities of hazardous materials

Substitute: replace a hazardous material by a less hazardous one

Attenuate (moderate): use the hazardous material in a less hazardous form

Simplify: avoid unnecessary complexity in facilities and processes, in order to make operating errors less likely.

Full inherent safety, i.e. total absence of hazards, is seldom if ever achievable. Therefore, it is well advised to avoid the absolute term “inherently safe” and instead refer to “inherently safer” technologies and procedures. For instance, it may be impossible to eliminate an explosive reactant. Usually, it is nevertheless possible to substantially reduce the hazard it gives rise to by drastically reducing the inventories of the substance. One way to do this is to produce the substance locally in a continuous process. In terms of the four above-mentioned strategies, this means that minimization is chosen instead of substitution.

The disaster in a chemical factory in Bhopal, India, in 1984, illustrates this. With an official death toll of 2259, it is the largest accident in the history of the chemical industry, and it has also been called “the worst example of an inherently unsafe design” (Edwards 2005, p. 91). Methyl isocyanate, the substance that caused the calamity, was an intermediate that was stored in large quantities (ibid.). The final product could have been obtained from the same raw materials via an alternative chain of reactions in which methyl isocyanate is not produced. This and other alternative processes should have been considered. Even if a process involving methyl isocyanate was chosen, storage of large quantities of the substance could and should have been avoided.

In general, solving a problem with inherent safety is preferable to relying on interventions at later stages in a potential chain of events leading up to an accident. A major reason for this is that as long as a hazard still exists, it can be activated by some unanticipated triggering event. Even with the best of control measures, some unforeseen event can give rise to an accident. Even if a dangerous material is safely contained in the ordinary process, there is always a risk that it will escape, for instance, due to a fire, an uncontrolled chemical reaction, sabotage or an unusual mistake (Hansson 2010). Even the best add-on safety technology can fail or be destroyed in the course of an accident. An additional reason is that inherent safety is usually more efficient against security threats than fail-safety. Add-on safety measures, which are typically required for fail-safety, can often easily be deactivated by those who wish to do so. When terrorists enter the plant with the intent to blow it up, it does not matter much if all ignition sources have been removed from the vicinity of explosive materials. They will bring their own ignition source. Similarly, even if a toxic substance has been securely contained in a closed process, they can usually find ways to release it. In contrast, if explosive and toxic substances have been removed or their quantities drastically reduced, then the plant is safer, not only against accidents but also against wilfully created disasters.

Safety measures based on the ideas of inherent safety have contributed much to reducing hazards in the chemical industry (Hendershot 1997; Overton and King 2006). However, several commentators have complained that progress in the implementation of inherent safety is too slow (Kletz 2004; Edwards 2005; Srinivasan and Natarajan 2012). Indeed, 24 years after the Bhopal accident, investigations of a fatal accident at a chemical plant in West Virginia revealed considerable safety problems in the plant. A large inventory of methyl isocyanate, up to 90,000 kg, was stored on the plant. Luckily, no detectable release of the substance took place in the 2008 accident. (The death toll of the Bhopal accident was due to release of 47,000 kg of the same substance.) Inherently safer alternatives to this massive storage of the substance had previously been considered but had been rejected as too expensive (Ogle et al. 2015).

It has often been proposed that the ideas of inherent safety should be exported to other industries, including mining, construction and transportation (Gupta and Edwards 2003). However, the only other industry in which inherent safety has a major role is the nuclear industry. Much effort has been devoted to developing nuclear reactors that are inherently safer than those currently in use. By this is meant that even in the case of failure of all active cooling systems and complete loss of coolant, fuel element temperatures should not exceed the limits below which most radioactive fission products remain confined within the fuel elements (Elsheikh 2013; Adamov et al. 2015).

Several authors have discussed the application of inherent safety to the construction of road vehicles and infrastructure. Inherent safety is often mentioned as a means to make progress towards Vision Zero for traffic safety. One recurrent idea is that speeds should be kept at levels low enough for the inherent safety of the system to prevent serious accidents (Tingvall and Haworth 1999; Khorasani-Zavareh 2011; Hakkert and Gitelman 2014). Arguably, much ongoing work in the construction of safer road vehicles can be described as applications of the basic principles of inherent safety. However, contrary to the literature on chemical and nuclear engineering, the technical literature on vehicle safety seldom refers to the notion of inherent safety.

The Substitution Principle

As we saw in the previous subsection, the substitution of hazardous substances by less dangerous ones is one of the major methods to achieve inherent safety. Independently of the inherent safety principle, a “substitution principle” has gained prominence in chemicals policy. The substitution principle requires the replacement of toxic chemicals by less dangerous alternatives. According to most versions of the principle, the replacement may be either another chemical or some non-chemical method to achieve the same or a similar result. The earliest example on record of a general rule requiring such substitutions seems to be a paragraph in the Swedish law on workplace health and safety from 1949:

A poisonous or otherwise noxious substance shall be replaced by a non-toxic or less harmful one whenever this can reasonably be done considering the circumstances. (Svensk författningssamling 1949, p. 401)

A special “substitution principle” for hazardous chemicals was introduced into the European health and safety legislation in 1989 (European Union 1989). It states that the employer has to implement preventive measures according to a series of “general principles of prevention,” one of which is “replacing the dangerous by the non-dangerous or the less dangerous” (European Union 1989, II.6.2). Substitution was also emphasized as a major risk-reducing strategy in the discussions in the 1990s that led up to a new European chemicals legislation (Sørensen and Petersen 1991; Antonsson 1995). The European Commission’s 2001 White Paper recommended “the substitution of dangerous by less dangerous substances where suitable alternatives are available” (European Commission 2001). Following this, a substitution principle was integrated into the European chemicals legislation (the REACH legislation), which was adopted in 2006. The substitution principle has also had an important role in various projects for chemical safety promoted by both government agencies and industrial companies (Lissner and Romano 2011; Hansson et al. 2011). Increasingly, the substitution principle has become associated with the movement for green chemistry, i.e. chemical engineering devoted to developing less hazardous chemical products and processes (Fantke et al. 2015; Tickner et al. 2019).

Decisions based on the substitution principle are often hampered by lack of reliable knowledge on the effects of both the chemicals currently in use and their potential alternatives (Rudén and Hansson 2010). Due to incomplete or inaccurate information, attempts to apply the substitution principle have sometimes led to the replacement of an unsafe product by another product that is in fact no better:

The chemical trichloroethylene (TCE), a volatile organic chemical, was widely used as a degreaser in the manufacture of electronic circuits and components until concerns about TCE’s environmental effects led the industry to replace it with trichloroethane (TCA), which has similar chemical structure. TCE and TCA were among the most widely used industrial degreasers, and they are now found in many of the hazardous cleanup sites listed on the National Priorities List. TCA, in turn, was replaced as a degreaser by chlorofluorocarbons such as Freon when ozone depletion concerns were raised about TCA in the 1990s. The use of Freon as a chemical degreaser was eventually phased out due to its own health and environmental concerns. Now, new mixtures of solvents are being used in vapor degreasing. (Bent 2012, pp. 1402–1403)

Generally speaking, the difficulties in assessing health risks and environmental risks are larger for chemical substances than for most other sources of potential hazards (Rudén and Hansson 2010). Therefore, a double strategy for chemical safety is advisable: Systematic work to replace hazardous substances and processes by less hazardous alternatives needs to be combined with equally methodical endeavours to reduce emissions and exposures.

In applications of the substitution principle, priority is usually given to substituting the most hazardous products and processes, but there is no predetermined level of risk below which further substitutions to even less perilous substances and methods are considered unnecessary. Furthermore, the principle is not “subordinated to purely economical considerations” (Szyszczak 1992, p. 10). This is in line with

the above-mentioned European health and safety legislation from 1989, which says the following:

The employer shall be alert to the need to adjust these measures to take account of changing circumstances and aim to improve existing situations. (European Union 1989, II.6.1)

Notably, this requirement is not restricted to companies in which existing conditions are below a certain standard. With this interpretation, the substitution principle is well in line with improvement principles such as Vision Zero and continuous improvement. It can also, with this interpretation, be classified as an improvement principle (Hansson 2019).

However, the substitution principle has sometimes been interpreted in ways that weaken its effects. In particular, high demands on the functionality of the replacement can sometimes block health and safety improvements. For instance, the substitution principle has been defined as “the replacement of a substance, process, product, or service by another that maintains the same functionality” (UK Chemicals Stakeholder Forum 2010). This would mean that a substitution can only be required if the replacement functions at least as well as the harmful substance that one wishes to avoid. To mention just one example, it would imply that a company using a highly toxic metal degreaser could not be required to substitute it by something less dangerous if the best replacement would require a small increase in the time that the metal parts have to be immersed in the solvent. With such an interpretation, the substitution principle would lose much of its effect (Hansson et al. 2011).

Safety Factors

A *safety factor* is a numerical factor (i.e. a number) that is used as a rule of thumb to create a margin to dangerous conditions. The most common uses of safety factors are in structural mechanics and in toxicology. In structural mechanics, to apply a safety factor x means to make a component x times stronger than what the predicted load requires. In toxicology, to apply a safety factor x means to only allow exposures that are at least x times smaller than some dose believed to be barely safe.

Safety factors provide a safety reserve, i.e. a distance or difference between the actual conditions and the conditions expected to cause a failure. You introduce a safety reserve if you hang your child’s swing with a stronger rope than what you actually believe to be necessary to hold a person using the swing. If you do this intuitively, the safety reserve is non-quantitative. If you ask the shop attendant for a rope that holds three times the highest load you expect, then your safety reserve is quantitative and expressible as a safety factor of three.

Non-quantitative safety reserves have been used in the building trades since prehistoric times (Randall 1976; Kurrer 2018). The early history of safety factors does not seem to have been written before, and a brief account will therefore be

given here. The earliest record of a quantitative safety factor may be a letter written in March 1812 by the English inventor Richard Trevithick (1771–1833), where he described how he used what we would today call a safety factor of 4 in the testing of steam engines:

To prevent mischief from bad castings, or from the fire injuring the surface of cast iron, I make the boilers of wrought iron, and always prove them with a pressure of water, forced in equal to four times the strength of steam intended to be worked with. (Trevithick 1872, p. 14)

The use of a safety factor for steam pressure seems to have been a common practice in Britain in the early nineteenth century. In his book on steam engines from 1822, the British science writer Charles Frederick Partington referred to four engine makers who all recommended the practice. However, they had widely different views on what an appropriate safety factor should be. One said that steam engines should be tested at 2 to 3 times higher pressure than the intended work pressure, another recommended 10 to 12 times higher pressure, a third 14 to 20 times higher, and a fourth 50 times higher (Partington 1822, pp. 109, 112, 113, and Appendix, p. 76). We can conclude that the notion of a safety factor was well known among engine makers at this time, although they neither had a name for it nor a common view on its value.

In his 1827 book on steam engines, the influential English civil engineer Thomas Tredgold (1788–1829) referred to the “excess of strength” that is required in a boiler. Although he wrote only five years later than Partington, he reported a consensus in the matter: “it has been almost universally allowed, that three times the pressure on the valve in the working state, should be borne by the boiler without injury.” However, he was critical of that consensus. He proposed that the factor of 3 could be lowered to 2 for “ordinary low-pressure steam boilers,” whereas high-pressure boilers required higher factors, depending on their construction (Tredgold 1827, pp. 257–258). His statement that a factor of 3 was customary is confirmed in a call for tenders for new locomotive steam engines that was sent out in 1830 by an American railway company. They stated that they considered themselves at liberty to put the engine “to the test of a pressure of water, not exceeding three times the pressure of the steam intended to be worked, without being answerable for any damage the machine may receive in consequence of such test” (Thomas 1830, p. 373).

In his three-volume book on bridge-building, published in 1850, the English civil engineer Edwin Clark (1814–1894) reproduced a text from 1846 describing the construction of a bridge such that “its breaking-weight is seven times as great as any weight with which in practice it can ever be loaded.” He called this number a “factor of safety” and discussed how it should be used in calculations, given that the bridge’s own weight had to be taken into account (Clark 1850, pp. 514–515). He reported that the famous Scottish engineer Robert Stephenson (1772–1850) favoured a factor of 7. This gives the impression that the use of safety factors was well established at the time, not only in boilermaking but also in civil engineering. Its earlier background in civil engineering remains to be investigated.

In 1859, the Scottish engineer and physicist William Rankine (1820–1872) published a table of “factors of safety” for different materials. This factor was, essentially, a ratio between breaking load and working load. He recommended safety factors of 10 for timber, 8 for stones and bricks and between 4 and 8 for different types of iron and steel (Rankine 1859, p. 65).

In 1873, the American engineer Barnet Le Van wrote a report to the Franklin Institute on a boiler explosion in Pennsylvania that had killed 13 persons and wounded many more. He concluded that proper maintenance and regular examination and testing of the boiler, in accordance with well-established routines, would have prevented the accident. However, he also had a more general conclusion:

In conclusion, I would call the attention of the Institute to the factor of safety for boilers as being entirely too low. The great number of disastrous explosions that have lately occurred in different parts of the country are the best evidences of the fact. The Bridge Engineers have long since come to this conclusion, and have fixed their factor of safety at one-eighth the ultimate value of the material. (Le Van 1873, p. 253)

The value of the safety factor for boilers that he criticized is not mentioned in his text, but it may well have been 3.

Today, safety factors are almost ubiquitous in engineering design. It is generally agreed that their main purpose is to compensate for five major sources of error in design calculations (Knoll 1976; Moses 1997):

1. Higher loads than those foreseen
2. Worse properties of the material than foreseen
3. Imperfect theory of the failure mechanism in question
4. Possibly unknown failure mechanisms
5. Human error (e.g. in design)

In toxicology, the first proposal to apply safety factors seems to have been Lehman’s and Fitzhugh’s proposal in 1954 to calculate the Acceptable Daily Intake of food additives by dividing the highest dose (in milligrams per kilo body weight) at which no effect had been observed in animals by 100 (Dourson and Stara 1983). Today, safety factors are essential components of regulatory food toxicology. They are also widely used in ecotoxicology. In both these applications, it is common to construct an overall safety factor by multiplying several safety factors for various uncertainties and variabilities. Thus, the traditional 100-fold factor is commonly accounted for as a combination of a factor of 10 for interspecies variability (between experimental animals and humans) in response to toxicity and another factor of 10 for intraspecies variability (among humans). In more recent approaches, toxicological safety factors often incorporate additional subfactors, referring, for instance, to differences between experimental and real-life routes of exposure, extrapolation from short-term experimental to life-long real-life exposures, and deficiencies in the available data (Gaylor and Kodell 2000). However, consistent use of safety factors has not been introduced into the process of setting occupational exposure limits. That area is

still dominated by case-by-case compromises between health protection and economic considerations, often resulting in exposure limits at levels where negative health effects are expected (see section “Exposure Limits”).

Since the 1990s, the use of safety factors in both structural engineering and toxicology has been criticized by scientists who want to replace them by calculated failure probabilities. However, in practice, the safety factor approach is still dominant, and it has only rarely been replaced by probability calculations. One reason for this is that probabilistic calculations are often much more complicated and time-consuming than the use of safety factors. Another reason is that meaningful probabilities are not available for some of the potential failures that safety factors are intended to protect against. In structural mechanics, this applies, for instance, to unknown failure mechanisms and imperfections in the calculations. In toxicology, it applies to unknown metabolic differences between species and unknown effects only occurring in parts of the human population (Doorn and Hansson 2011).

Multiple Safety Barriers

When several measures are employed to improve safety, they can often be perceived as a chain of safety measures or as they are then often called: a chain of safety barriers. Each of these barriers should be as independent as possible of its predecessors in the sequence, so that if the first barrier fails, then the second is still intact, etc. The use of multiple barriers is often advisable even if the first barrier is strong enough to withstand all foreseeable strains and stresses. The reason for this is that we cannot foresee everything. If the first barrier fails for some unforeseen reason, then the second barrier can provide protection.

The archetype of multiple safety barriers is an ancient fortress. If the enemy manages to pass the first wall, then there are additional layers that protect the defending forces. This is an age-old practice. As early as 3200 BCE, the Sumerian town Habuba Kabira (now in Syria) was surrounded by double walls (Keeley et al. 2007, p. 86). Double and triple walls were also erected around other major cities in the ancient Near East (Mielke 2012, p. 76). In the early Iron Age (around 450 BCE), hill forts were built in Britain with up to four concentric ramparts (Armit 2007).

Some engineering safety barriers exhibit the same spatial pattern as the concentric barriers of a fortification. Illustrative examples of this can be found in nuclear waste management. For instance, the nuclear industry in Sweden has proposed that spent fuel from the country’s nuclear reactors should be placed in copper canisters constructed to resist all foreseeable stresses. The canisters will be surrounded by a layer of bentonite clay, intended to protect against movements in the rock and to absorb radionuclides, should they leak from the canisters. This whole construction is placed in deep rock, in a geological formation that has been selected to minimize transportation to the surface of any possible leakage of radionuclides. The idea behind this construction is that the whole system of barriers should have a high degree of redundancy, so that if one of the barriers fails, then the remaining ones will

suffice to keep the radionuclides below the surface (Jensen 2017; Lersow and Waggitt 2020, pp. 282–287).

More generally, the safety measures (“barriers”) included in a multiple-barrier system should be arranged in a temporal or functional sequence, such that the second barrier is put to work if the first one fails, etc. The barriers may, but need not, be sequentially arranged in space. The combination of inherent safety and fail-safety can be used as an example of a temporally but not spatially sequential arrangement of barriers. Inherent safety is the first barrier. If it fails, then fail-safety should come in as a second resort. A systematic theoretical discussion of consecutive barriers in safety management was provided by William Haddon (1926–1985). He proposed that the protection against mechanical accidents such as traffic accidents should be conceptualized in terms of four types of barriers:

In the context of the recognition that abnormal energy exchanges are the fundamental cause of injury, accident prevention and hence accident research aimed at prevention are easily sorted into several types, each concerned with successive parts of the progression of events which lead up to these traumatic exchanges. In general, measures directed against accidental or deliberately inflicted injuries attempt: *first*, to prevent the marshalling of the hazardous energy itself, and *second*, if this is not feasible, to prevent or modify its release. *Third*, if neither of these is successful, they attempt to remove man from the vicinity, and *fourth*, if all of these fail, an attempt is made to interpose an appropriate barrier which will block or at least ameliorate its action on man. (Haddon 1963, p. 637)

For another example, consider a chemical process in which hydrogen sulfide is used as a raw material in the production of organosulfur compounds. Hydrogen sulfide is a deadly and treacherous gas, and it is therefore imperative to protect workers against exposure to it. This can be done with the help of a series of five barriers. The first barrier consists in reducing the use of the substance as far as possible. If it cannot be dispensed with completely, then resort must be had to the second barrier, which consists in encapsulating the process efficiently so that leakage of hydrogen sulfide is excluded as far as possible. The third barrier is careful maintenance, including regular checking of vulnerable details such as valves. The fourth barrier is an automatic gas alarm, combined with routines for evacuation of the premises in the case of an alarm. The fifth barrier is efficient and well-trained rescue and medical services. Importantly, even if the first, second, third and fourth of these barriers have been meticulously implemented, the fifth barrier should not be omitted. Doing so amounts to what we can call the “Titanic mistake.”

The sinking of the Titanic on April 15, 1912, is one of the most infamous technological failures in modern history. The ship was built with a double-bottomed hull that was divided into sixteen compartments, each constructed to be watertight. At least two of these could be filled with water without danger. Therefore, the ship was believed to be virtually unsinkable, and consequently, it was equipped with lifeboats only for about half of the around 2200 persons on-board. This was in line with the regulations at the time, which only required lifeboats for 990 persons for this ship (Hutchinson and de Kerbrech 2011, p. 112). Archibald Campbell Holms

(1861–1954), a prominent Scottish shipbuilder (also known as a leading spiritualist), commented as follows on the accident in his textbook on shipbuilding:

As showing the safety of the Atlantic passenger trade, may be pointed out that, of the six million passengers who crossed in the ten years ending June 1911, there was only a loss of six lives. The fact that *Titanic* carried boats for little more than half the people on board was not a deliberate oversight, but was in accordance with a deliberate policy that, when the subdivision of a vessel into watertight compartments exceeds what is considered necessary to ensure that she shall remain afloat after the worst conceivable accident, the need for lifeboats practically ceases to exist, and consequently a large number may be dispensed with. The fact that four or five compartments were torn open in *Titanic*, although no longer an inconceivable accident, may be regarded as an occurrence too phenomenal to be used wisely as a precedent in deciding the design and equipment of all passenger vessels in the future. (Holms 1917, p. 374)

Needless to say, this is an unusually clear example of the type of thinking that the concept of multiple safety barriers is intended to overcome. Luckily, most reactions to the accident were wiser than that of Campbell Holms. In consequence of the disaster, maritime regulations for long sea voyages were changed to require lifeboats for all passengers. However, the changes did not apply to shorter sea voyages. As late as in the 1960s, a night ferry between Belfast and the English seaport Heysham took up to 1800 passengers but had lifeboats only for 990 (Garrett 2007).

Redundancy

The notion of multiple barriers can be generalized to that of *redundancy*. By redundancy is meant that safety is upheld by a set of components or processes, such that more than one of them have to fail for conditions to become unsafe (Downer 2011; Hammer 1980, pp. 71–75). The redundant components can be arranged in different ways, for instance, in parallel or consecutively. If the arrangement is consecutive, then we have the special case of multiple barriers. Redundancy with a parallel arrangement can be exemplified by the engine redundancy in aircraft. This means that an airplane can reach its destination or at least the nearest airport, even if not all the engines are operative (DeSantis 2013). As this example shows, redundancy can be a way to achieve fail-safety.

The major difficulty in constructing redundant systems is to make the redundant parts as independent of each other as possible. If two or more of them are sensitive to the same type of impact, then one and the same destructive force can get rid of them in one fell swoop. For instance, any number of concentric walls around a fortified city could not protect the inhabitants against starvation under siege. Similarly, ten independent emergency lights in a tunnel can all be destroyed in a fire, or they may all be incapacitated due to the same mistake by the maintenance department. The Fukushima Daiichi nuclear accident in 2011 was caused by a natural disaster (an earthquake and its resultant tsunami), which shut down both the reactors' normal electricity supply and the emergency diesel generators. In consequence, the emergency cooling system did not work, which led to nuclear meltdowns and the release of radioactive material. This

would not have happened if the emergency generators had been placed at a higher altitude than the reactors. In general, how much safety is obtained with an arrangement for redundancy depends to a large degree on how sensitive the system is to failures affecting several redundant parts at the same time (“common-cause failures”). Often, safety is better served by few but independent barriers than by many barriers that are sensitive to the same sources of incapacitation.

The quality of redundancy systems is often discussed in terms of diversity and segregation. By *diversity* is meant that redundant parts differ in their constructions and mechanisms. For instance, in order to avoid dangerously high temperatures in a chemical reactor, we may introduce two temperature guards, each of which automatically turns off the reactor if a certain temperature limit is exceeded. The redundancy obtained by having two instruments is improved if they are of different types. It is also improved if we employ different software for their operations (Vilkomir and Kharchenko 2012). By *segregation* is meant that redundant components are physically separated from each other. This is done in order to reduce the risk of spatially limited common-cause failures produced, for instance, by fire, explosion, flooding, structural failure or sabotage. Segregation is more easily achieved in large industrial buildings or complexes than in operations with limited space such as ships, offshore platforms and aircrafts. However, the principle has been applied with success in the latter types of workplaces as well (Kim et al. 2017).

Summary

The various error tolerance principles that we have discussed in this section – fail-safety, inherent safety, substitution, safety factors, multiple safety barriers and redundancy – are all perfectly compatible with Vision Zero and other improvement principles. At least one of them, namely, inherent safety, has also been discussed in connection with Vision Zero. The error tolerance principles can all be seen as means to implement the improvement principles. In general, it is advisable to combine several error tolerance principles, as explained above in the subsections on multiple barriers and redundancy.

Evidence Evaluation Principles

Decisions on safety often have to be based on information that may be difficult to obtain. We may have to ask questions such as: Can this structure sustain the additional load we intend to place on it? Is this chemical exposure hazardous to human health? How reliable is the gas alarm? Sometimes, trustworthy answers to such questions can be obtained, but on other occasions, we have to make decisions based on uncertain or insufficient evidence. This section is devoted to such principles. We will begin with the *precautionary principle* and then discuss three of its alternatives, namely, *reversed burden of proof*, *risk neutrality* and “*sound science*”.

The Precautionary Principle

According to a common misconception, the precautionary principle says that all our decisions should be cautious. According to that reading of the principle, we all apply the precautionary principle when we wear a seat belt or have our children vaccinated. But this is not what the precautionary principle means. It is a well-defined principle for the evaluation of evidence, defined in international treaties and also in the European legislation. What it means is, essentially, that even if the evidence of a danger is uncertain, we may, and often should, take precautionary measures against it.

This is of course no new way of thinking. Presumably, our ancestors refrained from entering a cave if they heard a suspicious growl from it, even if they were far from convinced that the animal they heard was dangerous. An illustrative, more recent example is the closing of a water pump in London in 1854. In early September that year, the city was struck by cholera, and 500 people died in 10 days. The physician John Snow notified the authorities that according to his investigations, a large number of those affected by the disease had drunk water from a pump on Broad Street. The authorities had no means to verify that this was more than a coincidence. According to the prevalent opinion among physicians, cholera was transmitted through air rather than water. However, although the evidence was uncertain, the authorities decided to have the handle removed from the pump. This had the effect hoped for, and the cholera epidemic was curbed (Snow 2002; Koch and Denike 2009).

The modern precautionary principle had precursors in Swedish and German legislation and in treaties on protection of the North Sea in the 1980s (Hansson 2018b). It rose to international importance through the Rio Declaration on Environment and Development that was a major outcome of the 1992 so-called Earth Summit in Rio de Janeiro:

Principle 15. Precautionary principle

In order to protect the environment, the precautionary approach shall be widely applied by States according to their capabilities. Where there are threats of serious or irreversible damage, lack of full scientific certainty shall not be used as a reason for postponing cost-effective measures to prevent environmental degradation. (United Nations 1992)

The European Union and several of its member states have incorporated the precautionary principle into their legislations. Through the 1992 Maastricht amendments to the European Treaty (Treaty of Rome, now known as the Treaty on the Functioning of the European Union), the precautionary principle was written into European legislation (Pyhälä et al. 2010, p. 206). According to the treaty, union policy “shall be based on the precautionary principle” (European Union 2012). In February 2000, the Commission issued a Communication on the Precautionary Principle that further clarified its meaning:

The precautionary principle is not defined in the Treaty, which prescribes it only once – to protect the environment. But *in practice*, its scope is much wider, and specifically where preliminary objective scientific evaluation, indicates that there are reasonable grounds for

concern that the potentially dangerous effects on the *environment, human, animal or plant health* may be inconsistent with the high level of protection chosen for the Community. . .

Recourse to the precautionary principle presupposes that potentially dangerous effects deriving from a phenomenon, product or process have been identified, and that scientific evaluation does not allow the risk to be determined with sufficient certainty.

The implementation of an approach based on the precautionary principle should start with a scientific evaluation, as complete as possible, and where possible, identifying at each stage the degree of scientific uncertainty. . .

[M]easures based on the precautionary principle should be maintained so long as scientific information is incomplete or inconclusive, and the risk is still considered too high to be imposed on society, in view of chosen level of protection. (European Commission 2000)

It can clearly be seen from this and other official texts that the precautionary principle is a principle for decision-making in situations with uncertainty about a potential hazard. In the United States, the precautionary principle has seldom been invoked by policymakers, but some administrations have taken measures based on similar thinking, using other terms such as “better safe than sorry” (Wiener and Rogers 2002). Even in Europe, the precautionary principle is seldom referred to outside of the policy areas that concern health, safety and the environment. However, similar approaches to evidence are prevalent in a wide range of policy areas although no reference is made to the precautionary principle. For instance, economists commonly agree that action should be taken against a potential financial crisis even in the absence of full evidence that it will otherwise take place. Similarly, a military commander who waits for full evidence of an enemy attack before taking any countermeasures would be regarded as incompetent.

The following two, hypothetical but realistic, examples can be used to exemplify the precautionary principle:

The volcano example

A group of children are tenting close to the top of an old volcano that has not been active for thousands of years. While they are there, seismographs and gas detectors suddenly indicate that a major eruption may be on its way. A committee of respected volcanologists immediately convene to evaluate the findings. They conclude that the evidence is uncertain but weighs somewhat in the direction that a major eruption will take place in the next few days. They unanimously conclude that although the evidence is not conclusive, it is more probable that an eruption is imminent than that it is not. (Hansson 2018b, pp. 269)

The baby food example

New scientific evidence indicates that a common preservative agent in baby food may have a small negative effect on the child’s brain development. According to the best available scientific expertise, the question is far from settled but the evidence weighs somewhat in the direction of there being such an effect. A committee of respected scientists unanimously concluded that although the evidence is not conclusive, it is more probable that the effect exists than that it does not. The food safety agency has received a petition whose signatories request the immediate prohibition of the preservative. (Hansson 2018b, pp. 268–269)

As these examples exemplify, there are occasions when we wish to take measures against a possible danger, although the scientific information is not sufficient to

establish that the danger is real. At least in the second case, this would (in Europe) be described as an application of the precautionary principle.

However, there are also dangers with basing decisions on less than full scientific evidence. If we give up the scientific basis entirely, then we run the risk of making decisions that have no foundation at all, leaving room for decisions based on prejudice and uninformed suppositions (Hansson 2016, 2018a). It is necessary to ensure that full use is made of the available scientific information even when we are willing to base decisions on incomplete evidence. The following three principles have been proposed as guidelines (Hansson 2008, pp. 145–146). They can be called the principles of *science-based precaution* in practical decision-making:

1. The evidence taken into account in the policy process should be the same as in a purely scientific evaluation of the issue at hand. Policy decisions are not well served by the use of irrelevant data or the exclusion of relevant data.
2. The assessment of how strong the evidence is should be the same in the two processes.
3. The two processes may differ in the *required* level of evidence. It is a policy issue how much evidence is needed for various practical decisions.

A Reversed Burden of Proof

In particular in the discussion on chemical hazards, it has often been claimed that the onus of proof should fall to those who claim that a substance can be used without danger, rather than those who wish to restrict its use. This is commonly called the “reversed burden of proof” (Wahlström 1999, pp. 60–61). If by the burden of proof is meant the duty to pay for the required investigations of the effects of a substance, then this is a burden that can and arguably should be borne by those who wish to put the substance on the market. In many jurisdictions, considerable duties of investigation have already been imposed on companies wishing to put a chemical substance on the market. However, in discussions of chemical risks, the term “burden of proof” usually means something else, which is close to what legal scholars call “burden or persuasion”: It is claimed that unless the company in question can prove that the substance is harmless, the substance should not be used. On the face of it, this seems to be an excellent safety principle: We should only use provenly harmless substances. Who can be against that?

Unfortunately, this principle has a fundamental defect: It cannot be realized. It can often be proved beyond reasonable doubt that a substance has a particular adverse effect. However, it is often impossible to prove beyond reasonable doubt that a substance does not have a particular adverse effect, and in practice, it is always impossible to prove that it has no adverse effect at all (Hansson 1997a). The major reason for this is that with respect to serious health effects, we care about risks that are small in comparison to the limits of detection in scientific studies. If we only

cared about whether an exposure kills more than one-tenth of the exposed population, then this problem would not arise. But for ethical reasons, we wish to exclude even much lower frequencies of adverse effects.

As a rough rule of thumb, epidemiological studies can only detect reliably excess risks that are about a tenth of the risk in the unexposed population. For instance, suppose that the lifetime risk of a deadly heart attack (myocardial infarction) is 10% in a population. Furthermore, suppose that a part of the population is exposed to a substance that increases this risk to 11%. This is a considerable risk increase, leading to the death of one in a hundred of those exposed. However, even in large and well-conducted epidemiological studies, chances are slim of detecting such a difference between the exposed and the unexposed group (Vainio and Tomatis 1985). There are similar statistical problems in animal experiments (Weinberg 1972, p. 210; Freedman and Zeisel 1988; Hansson 1995).

The lesson from this is that it is in general impossible to prove that an exposure has no negative effects. Demands for such proofs can be counterproductive since they contribute to the misconception that chemical risks can be eliminated with substance choice, without any measures to reduce exposure. A realistic strategy to minimize chemical risks should be based on a multiple-barrier approach. Appropriate pre-market investigations of substances, constructed to discover negative health effects as far as possible, can serve as a first barrier. However, this has to be followed by other barriers, including measures that reduce exposures as well as check-ups to discover unexpected harmful effects.

Many technological devices are accessible to more reliable pre-market testing than chemical substances. The reason for this is that relevant failure types, such as mechanical and electrical failures, are much better understood than toxicity, which makes more reliable testing possible. (A caveat: This does not always apply to software failures.) However, this does not exclude the need for a “second barrier” in the form of post-marketing follow-ups. Experiences from safety recalls in the motor vehicle, aircraft, toy, food, pharmaceutical and medical device industries show that even in industries with a comparatively high focus on safety, the “first barrier” of pre-market testing and assessment does not always exclude the marketing of unsafe products (Rupp 2004; Bates et al. 2007; Berry and Stanek 2012; Nagaich and Sadhna 2015; Shang and Tonsor 2017; Niven et al. 2020; Johnston and Harris 2019). Proposals have been made to introduce routines for safety recalls in industries still lacking recall traditions, such as the building materials industry (Huh and Choi 2016; Bowers and Cohen 2018; Watson et al. 2019).

Risk Neutrality

Opponents of the precautionary principle have often proposed that it should be replaced by a risk-neutral or, in a common but rather misleading terminology, “risk-based” approach (Klinke et al. 2006, p. 377). By this is meant that risks should be assessed according to their expectation values, i.e. the product of some measure of

the expected damage with its probability. This is the way in which risks are assessed in cost-benefit analysis. As we saw above, this is a method with considerable drawbacks. Attempts to use it as a replacement for the precautionary principle will also run into an additional, quite severe problem: The precautionary principle is a principle for the interpretation of uncertain or limited evidence. For that task, meaningful probabilities are usually not available. In practice, “risk-based” decision-making tends to proceed by neglecting uncertainties and only taking known dangers into account.

The following, somewhat stylized, example serves to illustrate the point: Consider two substances A and B, both of which are alternatives for being used in an application where they will leak into the aquatic environment. A has been thoroughly tested and is known to be weakly ecotoxic. It is not known whether B is ecotoxic. (No exotoxicity was discovered in the standard tests, but due its chemical structure, some researchers have expressed worries that it may be toxic to other organisms than those included in those tests.) However, B is known to be highly persistent and bioaccumulative. This means that *if* B is ecotoxic, then it can be highly potent since it will accumulate in biota. The ecological risks of using substance A can be quantified and entered into a cost-benefit analysis and a “risk-based” decision procedure. However, since no meaningful probability can be assigned to the eventuality that B is ecotoxic, we cannot perform a “risk-based” assessment of its potential to harm the environment. Therefore, a “risk-based” assessment will show that A poses an ecological risk, but it will have no risk to report for B. In contrast, an assessment in line with the precautionary principle will put focus on the serious but unquantifiable risks that B may give rise to. Thus, in spite of its name, a “risk-based” assessment will in this case tend to downplay risks that are taken seriously if the precautionary principle is applied.

“Sound Science”

If “sound science” means good science, then all rational decision-makers should make use of sound science, combining it with decision criteria that are appropriate for the purposes of the decision. However, in recent discussions, the phrase “sound science” has acquired a different meaning. It was adopted as a political slogan in 1993, when the tobacco company Philip Morris initiated and funded an ostensibly independent organization called The Advancement of Sound Science Coalition (TASSC). Its major task was to promulgate pseudoscience in support of the claim that the evidence for health risks from passive smoking was insufficient for regulatory action (Ong and Glantz 2001). The term “sound science” has also been used in similar lobbying activities against reductions in human exposure to other toxic substances (Rudén and Hansson 2008, pp. 300–301; Samet and Burke 2001; Francis et al. 2006). Considerable efforts have been made to create “sound science” alternatives to the scientific consensus on climate change summarized by the IPCC (Cushman 1998; Boykoff 2007, p. 481; Dunlap and McCright 2010, p. 249; Hansson 2017).

The major effect of the requirements for “sound science” has been to delay and prevent health, safety and environmental regulations by incessantly questioning the evidence on which they are based (Neff and Goldman 2005). Decision-making based on uncertain evidence is consistently repudiated. However, disregarding well-grounded evidence of danger whenever it is not strong enough to dispel all doubts is nothing less than blatantly irrational. Even if you do not know for sure that a dog bites, reasonable suspicions that it does are reason enough to prevent your child from playing with the dog. Scientific evidence of danger should be treated in the same way.

Summary

In this section, we have studied four approaches to the evaluation of uncertain evidence. The precautionary principle, interpreted in the science-based way described above, is fully compatible with improvement principles such as Vision Zero, and it can be used to support their implementation. The idea of a reversed burden of proof, in its most common interpretation, is much more problematic. It cannot be implemented in practice, and its promotion tends to support a once-and-for-all approach to chemical safety, rather than a more appropriate multiple-barrier approach. Risk-neutral (“risk-based”) assessments of uncertain evidence are usually not feasible since they require probability values that cannot be obtained. Finally, “sound science,” in the sense that the phrase has acquired through the activities of tobacco lobbyists and their allies, should not be classified as a safety principle. It epitomizes the kind of risk-taking that has always stood in the way of safety.

Conclusion

We began our exploration of safety principles with an overview of how zero goals and targets have been used in widely different areas. We found that strivings for zero of something undesirable have the important advantage of counteracting fatalism and complacency. After that, we broadened our attention to a larger group of safety principles, containing continuous improvement, as low as reasonably achievable (ALARA) and best available technology (BAT). All these principles can be called improvement principles, since they convey the message that no level of risk above zero is fully satisfactory and that consequently, improvements in safety should always be striven for as long as they are at all possible. These principles are all fully compatible with each other, and we can see the different improvement principles as different ways to express the same basic message.

Next, we explored some other principles that tell us what levels of safety or risk reduction we should aim at (aspiration principles). Several of these principles tend to classify some unsafe and improvable conditions as acceptable. Such principles are not easily combined with Vision Zero and other improvement principles. However, we also found that one of these aspiration principles, namely, cost-effectiveness analysis,

fits in very well with the improvement principles. Cost-effectiveness analysis can be used to choose the safety measures that yield the largest improvements.

We then turned to the error tolerance principles, which are safety principles telling us that since failures are unavoidable, we have to ensure that the consequences of failures will be as small as possible. We discussed six such principles: fail-safety, inherent safety, substitution, safety factors, multiple safety barriers and redundancy. All of these principles are highly compatible with Vision Zero and other improvement principles. We can see them as complementary strategies for implementing the improvement principles.

Finally, we considered four evidence evaluation principles. One of them, namely, the precautionary principle, is well in line with Vision Zero and the other improvement principles.

Safety work is complex and in need of guidance on many levels. Therefore, we need several safety principles. Vision Zero and other improvement principles can tell us what we should aspire to. Error tolerance principles provide essential insights on the means that can lead us in that direction. We can use cost-effectiveness analysis to prioritize among the measures that are available to us and the precautionary principle to deal with uncertainties in the evidence available to us.

Cross-References

- ▶ [Suicide in the Transport System](#)
- ▶ [Vision Zero and Other Road Safety Targets](#)
- ▶ [Vision Zero in Disease Eradication](#)
- ▶ [Vision Zero in Suicide Prevention and Suicide Preventive Methods](#)
- ▶ [Vision Zero in Workplaces](#)
- ▶ [Vision Zero on Fire Safety](#)
- ▶ [What Is a Vision Zero Policy? Lessons from a Multi-sectoral Perspective](#)
- ▶ [Zero-Waste: A New Sustainability Paradigm for Addressing the Global Waste Problem](#)

Acknowledgement I would like to thank Karin Edvardsson Björnberg, Claes Tingvall and Henok Girma Abebe for valuable comments on an earlier version of this chapter.

References

- Abelson, P. H. (1988). Competitiveness: A long-enduring problem. *Science*, 240(4854), 865–865.
- Ackerman, B. A., & Stewart, R. B. (1988). Reforming environmental law: The democratic case for market incentives. *Columbia Journal of Environmental Law*, 13(2), 171–199.
- Adamov, E. O., Orlov, V. V., Rachkov, V. I., Slessarev, I. S., & Khomyakov, Y. S. (2015). Nuclear energy with inherent safety: Change of outdated paradigm, criteria. *Thermal Engineering*, 62(13), 917–927.

- American Psychological Association Zero Tolerance Task Force. (2008). Are zero tolerance policies effective in the schools? An evidentiary review and recommendations. *American Psychologist*, 63, 852–862.
- Anderson, E. (1988). Values, risks and market norms. *Philosophy and Public Affairs*, 17, 54–65.
- Anon. (1874). *The public general acts of the United Kingdom of Great Britain and Ireland passed in the thirty-seventh and thirty-eighth years of the reign of Her Majesty Queen Victoria*. London: George Edward Eyre and William Spottiswoode.
- Anon. (1902). Dead man's handle. *Electrical World and Engineer*, 39(7), 312–312.
- Anon. (1965). *A guide to zero defects. Quality and reliability assurance handbook* (4155.12-H). Washington, DC: Office of the Assistant Secretary of Defense. Downloaded December 2, 2019 from <https://apps.dtic.mil/dtic/tr/fulltext/u2/a950061.pdf>
- Anon. (1971). Zero in on safety. *Langley Researcher*, February 5 1971, p. 6.
- Anon. (2017). The legacy of zero tolerance policing. *New York Times*, February 20.
- Antonsson, A. B. (1995). Substitution of dangerous chemicals –The solution to problems with chemical health hazards in the work environment? *American Industrial Hygiene Association Journal*, 56(4), 394–397.
- Armit, I. (2007). Hillforts at war: From Maiden Castle to Taniwaha Pā. *Proceedings of the Prehistoric Society*, 73, 25–38.
- Auxier, J. A., & Dickson, H. W. (1983). Guest editorial: Concern over recent use of the ALARA philosophy. *Health Physics*, 44, 595–600.
- Baghel, N. B. A. (2005). An overview of continuous improvement: From the past to the present. *Management Decision*, 43(5), 761–771.
- Ball, J., & Procter, D. (1994). Zero-accident approach at British Steel, Teesside Works. *Total Quality Management*, 5(3), 97–104.
- Banatvala, J. E., & Brown, D. W. (2004). Rubella. *Lancet*, 363(9415), 1127–1137.
- Batalden, P. B., & Stoltz, P. K. (1993). A framework for the continual improvement of health care: Building and applying professional and improvement knowledge to test changes in daily work. *Joint Commission Journal on Quality Improvement*, 19(10), 432–452.
- Bates, H., Holweg, M., Lewis, M., & Oliver, N. (2007). Motor vehicle recalls: Trends, patterns and emerging issues. *Omega*, 35(2), 202–210.
- Belin, M.-Å., Tillgren, P., & Vedung, E. (2012). Vision Zero—a road safety policy innovation. *International Journal of Injury Control and Safety Promotion*, 19(2), 171–179.
- Benecke, O., & DeYoung, S. E. (2019). Anti-vaccine decision-making and measles resurgence in the United States. *Global Pediatric Health*, 6, 1–5.
- Bent, J. R. (2012). An incentive-based approach to regulating workplace chemicals. *Ohio State Law Journal*, 73, 1389–1455.
- Bergman, B. (2018). Quality principles and their applications to safety. In N. Möller, S. O. Hansson, J.-E. Holmberg, & C. Rollenhagen (Eds.), *Handbook of safety principles* (pp. 333–360). Hoboken: Wiley.
- Berry, M. G., & Stanek, J. J. (2012). The PIP mammary prosthesis: A product recall study. *Journal of Plastic, Reconstructive and Aesthetic Surgery*, 65(6), 697–704.
- Berwick, D. M. (1989). Continuous improvement as an ideal in health care. *New England Journal of Medicine*, 320(1), 53–56.
- Berwick, D. M. (2008). The science of improvement. *JAMA*, 299(10), 1182–1184.
- Berwick, D. M., Godfrey, A. B., & Roessner, J. (1990). *Curing health care. New strategies for quality improvement* (1st ed.). San Francisco: Jossey-Bass.
- Bessant, J., Caffyn, S., Gilbert, J., Harding, R., & Webb, S. (1994). Rediscovering continuous improvement. *Technovation*, 14(1), 17–29.
- Bicevskis, A. (1982). Unacceptability of acceptable risk. *Search*, 13(1–2), 31–34.
- Blomquist, G. C. (1988). *The regulation of motor vehicle and traffic safety*. Boston: Kluwer.
- Bollinger, R. E., Clark, D. G., Dowell, A. M., III, Ewbank, R. M., Hendershot, D. C., Lutz, W. K., Meszaros, S. I., Park, D. E., & Wixom, E. D. (1996). *Inherently safer chemical processes – A life*

- cycle approach*. New York: Center for Chemical Process Safety of the American Institute of Chemical Engineers.
- Bouton, E. M. (1924). Variable voltage control systems as applied to electric elevators. *Transactions of the American Institute of Electrical Engineers*, 43, 199–219.
- Bowers, S., & Cohen, D. (2018). How lobbying blocked European safety checks for dangerous medical implants. *BMJ*, 363, k4999.
- Bowling, B. (1999). The rise and fall of New York murder: Zero tolerance or crack's decline? *British Journal of Criminology*, 39(4), 531–554.
- Boykoff, M. T. (2007). From convergence to contention: United States mass media representations of anthropogenic climate change science. *Transactions of the Institute of British Geographers*, 32(4), 477–489.
- Broadbent, M. V., & Hubbard, L. B. (1992). Science and perception of radiation risk. *Radiographics*, 12, 381–392.
- Brues, A. M. (1958). Critique of the linear theory of carcinogenesis. *Science*, 128(3326), 693–699.
- Bush, B. J. (2012). Addressing the regulatory collapse behind the Deepwater horizon oil spill: Implementing a 'best available technology' regulatory regime for Deepwater oil exploration safety and cleanup technology. *Journal of Environmental Law and Litigation*, 26, 535–568.
- Butler, K. (2017). *Zero harm: Myth or reality*. PhD thesis, Queensland University of Technology.
- Cadier, B., Durand-Zaleski, I., Thomas, D., & Chevreul, K. (2016). Cost effectiveness of free access to smoking cessation treatment in France considering the economic burden of smoking-related diseases. *PLoS One*, 11(2), e0148750.
- Callaway, E. (2016). Dogs thwart effort to eradicate Guinea worm. *Nature*, 529(7584), 10–11.
- Cameron, C. (1975). Accident proneness. *Accident Analysis and Prevention*, 7, 49–53.
- Carlet, J., Fabry, J., Amalberti, R., & Degos, L. (2009). The 'zero risk' concept for hospital-acquired infections: A risky business! *Clinical Infectious Diseases*, 49(5), 747–749.
- Choksey, M. S., & Malik, I. A. (2004). Zero tolerance to shunt infections: Can it be achieved? *Journal of Neurology Neurosurgery & Psychiatry*, 75(1), 87–91.
- Clark, E. (1850). *The Britannia and Conway tubular bridges* (Vol. II). London: Day and Son.
- Cochi, S. L. (2017). Pivoting from polio eradication to measles and rubella elimination: A transition that makes sense both for children and immunization program improvement. *The Pan African Medical Journal*, 27(Suppl 3), 10.
- Coleman, K. (2005). *A history of chemical warfare*. Houndmills: Palgrave Macmillan.
- Cook, W. A. (1987). *Occupational exposure limits – Worldwide*. Akron: American Industrial Hygiene Association.
- Corten, O. (1999). The notion of 'reasonable' in international law: Legal discourse, reason and contradictions. *International and Comparative Law Quarterly*, 48, 613–625.
- Crosby, P. B. (1979). *Quality is free: The art of making quality certain*. New York: McGraw-Hill.
- Cushman, J. H. (1998). Industrial group plans to battle climate treaty. *New York Times*, April 26.
- Das, T. K. (2005). Zero discharge technology. In T. K. Das (Ed.), *Toward zero discharge: Innovative methodology and technologies for process pollution prevention* (pp. 225–261). Hoboken: Wiley-Interscience.
- Dekker, S. (2017). Zero commitment: Commentary on Zwetsloot et al., and Sherratt and Dainty. *Policy and Practice in Health and Safety*, 15(2), 124–130.
- Dekker, S. (2019). *Foundations of safety science: A century of understanding accidents and disasters*. Boca Raton: CRC Press.
- DeSantis, J. A. (2013). Engines turn or passengers swim: A case study of how ETOPS improved safety and economics in aviation. *Journal of Air Law and Commerce*, 77(4), 3–68.
- Diaz, S., et al. (2019). *Summary for policymakers of the global assessment report on biodiversity and ecosystem services of the Intergovernmental Science-Policy Platform on Biodiversity and Ecosystem Services*. Advance unedited version, 6 May 2019. Downloaded December 7, 2019 from https://ipbes.net/sites/default/files/downloads/spm_unedited_advance_for_posting_htn.pdf

- Doom, N., & Hansson, S. O. (2011). Should safety factors replace probabilistic design? *Philosophy and Technology*, 24, 151–168.
- Dourson, M. L., & Stara, J. F. (1983). Regulatory history and experimental support of uncertainty (safety) factors. *Regulatory Toxicology and Pharmacology*, 3(3), 224–238.
- Downer, J. (2011). On audits and airplanes: Redundancy and reliability-assessment in high technologies. *Accounting, Organizations and Society*, 36(4–5), 269–283.
- Dunlap, R. E., & McCright, A. M. (2010). Climate change denial: Sources, actors and strategies. In C. Lever-Tracy (Ed.), *Routledge handbook of climate change and society* (pp. 240–259). Abington: Routledge.
- Eberhard, M. L., Cleveland, C. A., Zirimwabagabo, H., Yabsley, M. J., Ouakou, P. T., & Ruiz-Tiben, E. (2016). Guinea worm (*Dracunculus medinensis*) infection in a wild-caught frog, Chad. *Emerging Infectious Diseases*, 22(11), 1961–1962.
- Edvardsson, K., & Hansson, S. O. (2005). When is a goal rational? *Social Choice and Welfare*, 24, 343–361.
- Edwards, D. W. (2005). Are we too risk-averse for inherent safety? An examination of current status and barriers to adoption. *Process Safety and Environmental Protection*, 83, 90–100.
- Elsheikh, B. M. (2013). Safety assessment of molten salt reactors in comparison with light water reactors. *Journal of Radiation Research and Applied Sciences*, 6(2), 63–67.
- Elvik, R. (1999). Can injury prevention efforts go too far? Reflections on some possible implications of Vision Zero for road accident fatalities. *Accident Analysis and Prevention*, 31(3), 265–286.
- Ensaldo-Carrasco, E., Andrew, C.-S., Cresswell, K., Bedi, R., & Sheikh, A. (2018). Developing agreement on never events in primary care dentistry: An international eDelphi study. *British Dental Journal*, 224(9), 733–740.
- European Commission. (2000). *Communication from the Commission of 2 February 2000 on the precautionary principle*, http://ec.europa.eu/dgs/health_consumer/library/pub/pub07_en.pdf
- European Commission. (2001). *Strategy for a future chemicals policy, White Paper*. European Commission 27.2.2001, <http://europa.eu.int/eurlex/en/com/wpr/2001/com20010088en01.pdf>
- European Union. (1989). Council Directive 89/391/EEC of 12 June 1989 on the introduction of measures to encourage improvements in the safety and health of workers at work. *Official Journal of the European Communities* L 183/1 29/6 1989.
- European Union. (2012). *Consolidated version of the treaty on the functioning of the European Union*, 26.10.2012. <http://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:12012E/TXT&from=EN>
- Evang, K. (1961). Summing up. *International Road Safety and Traffic Review*, 9(2), 40–42. & 50.
- Færdselssikkerhedskommissionen [Danish Road Safety Commission]. (2000). *Hver Ulykke Er en for Meget – Trafiksikkerhed Starter med Dig* [Every accident is one too many – Traffic safety starts with you]. Copenhagen: Trafikministeriet.
- Fantke, P., Weber, R., & Scheringer, M. (2015). From incremental to fundamental substitution in chemical alternatives assessment. *Sustainable Chemistry and Pharmacy*, 1, 1–8.
- Fenner, F. (1993). Smallpox: Emergence, global spread, and eradication. *History and Philosophy of the Life Sciences*, 15(3), 397–420.
- Fiksel, J. (1985). Toward a de minimis policy in risk regulation. *Risk Analysis*, 5, 257–259.
- Francis, J. A., Shea, A. K., & Samet, J. M. (2006). Challenging the epidemiologic evidence on passive smoking: Tactics of tobacco industry expert witnesses. *Tobacco Control*, 15(Suppl 4), iv68–iv76.
- Freedman, D. A., & Zeisel, H. (1988). From mouse-to-man: The quantitative assessment of cancer risks. *Statistical Science*, 3(1), 3–28.
- Furness, J. W. (1978). The motor vehicle – A good target for legislation. *Proceedings of the Institution of Mechanical Engineers*, 192(1), 157–170.
- Garg, S., & Singh, R. (2013). Getting to zero: Possibility or propaganda? *Indian Journal of Sexually Transmitted Diseases*, 34(1), 1–4.
- Garrett, D. J. (2007). Letter. *Engineering Management Journal*, 17(5), 8–8.

- Gaylor, D. W., & Kodell, R. L. (2000). Percentiles of the product of uncertainty factors for establishing probabilistic reference doses. *Risk Analysis*, 20(2), 245–250.
- Goh, T. N., & Xie, M. (1994). New approach to quality in a near-zero defect environment. *Total Quality Management*, 5(3), 3–10.
- Government Bill 1996/97:137. *Nollvisionen och det Trafiksäkra Samhället* [Vision Zero and the traffic safe society]. Stockholm: Government Offices of Sweden.
- Greenbaum, D. (2003). A historical perspective on the regulation of particles. *Journal of Toxicology and Environmental Health Part A*, 66(16–19), 1493–1498.
- Greene, J. A. (1999). Zero tolerance: A case study of police policies and practices in New York City. *Crime & Delinquency*, 45(2), 171–187.
- Gupta, J. P., & Edwards, D. W. (2003). A simple graphical method for measuring inherent safety. *Journal of Hazardous Materials*, 104, 15–30.
- Haddon, W., Jr. (1963). A note concerning accident theory and research with special reference to motor vehicle accidents. *Annals of the New York Academy of Sciences*, 107(2), 635–646.
- Hakkert, A. S., & Gitelman, V. (2014). Thinking about the history of road safety research: Past achievements and future challenges. *Transportation Research, Part F*, 25, 137–149.
- Halpin, J. F. (1966). *Zero defects – A new dimension in quality assurance*. New York: McGraw-Hill.
- Hammer, W. (1980). *Product safety management and engineering*. Englewood Cliffs: Prentice-Hall.
- Hampshire, S. (1972). *Morality and pessimism*. Cambridge: Cambridge University Press.
- Hansson, S. O. (1995). The detection level. *Regulatory Toxicology and Pharmacology*, 22, 103–109.
- Hansson, S. O. (1997a). Can we reverse the burden of proof? *Toxicology Letters*, 90, 223–228.
- Hansson, S. O. (1997b). Critical effects and exposure limits. *Risk Analysis*, 17, 227–236.
- Hansson, S. O. (1998a). *Setting the limit. Occupational health standards and the limits of science*. Oxford: Oxford University Press.
- Hansson, S. O. (1998b). A case study of pseudo-science in occupational medicine. *New Solutions*, 8, 175–189.
- Hansson, S. O. (2004a). Fallacies of risk. *Journal of Risk Research*, 7, 353–360.
- Hansson, S. O. (2004b). Weighing risks and benefits. *Topoi*, 23, 145–152.
- Hansson, S. O. (2007a). Ethics and radiation protection. *Journal of Radiological Protection*, 27, 147–156.
- Hansson, S. O. (2007b). Hypothetical retrospection. *Ethical Theory and Moral Practice*, 10, 145–157.
- Hansson, S. O. (2007c). Philosophical problems in cost-benefit analysis. *Economics and Philosophy*, 23, 163–183.
- Hansson, S. O. (2008). Regulating BFRs – From science to policy. *Chemosphere*, 73, 144–147.
- Hansson, S. O. (2010). Promoting inherent safety. *Process Safety and Environmental Protection*, 88, 168–172.
- Hansson, S. O. (2013a). *The ethics of risk. Ethical analysis in an uncertain world*. New York: Palgrave Macmillan.
- Hansson, S. O. (2013b). ALARA: What is reasonably achievable? In D. Oughton & S. O. Hansson (Eds.), *Social and ethical aspects of radiation risk management* (pp. 143–156). Amsterdam: Elsevier.
- Hansson, S. O. (2016). How to be cautious but open to learning: Time to update biotechnology and GMO legislation. *Risk Analysis*, 36(8), 1513–1517.
- Hansson, S. O. (2017). Science denial as a form of pseudoscience. *Studies in History and Philosophy of Science*, 63, 39–47.
- Hansson, S. O. (2018a). Risk, science and policy – A treacherous triangle. *Ethical Perspectives*, 25(3), 391–418.
- Hansson, S. O. (2018b). The precautionary principle. In N. Möller, S. O. Hansson, J.-E. Holmberg, & C. Rollenhagen (Eds.), *Handbook of safety principles* (pp. 258–283). Hoboken: Wiley.

- Hansson, S. O. (2019). Improvement principles. *Journal of Safety Research*, 69, 33–41.
- Hansson, S. O., & Schenk, L. (2016). Protection without discrimination: Pregnancy and occupational health regulations. *European Journal of Risk Regulation*, 7, 404–412.
- Hansson, S. O., Molander, L., & Rudén, C. (2011). The substitution principle. *Regulatory Toxicology and Pharmacology*, 59, 454–460.
- Harpur, N. F. (1958). Fail-safe structural design. *Aeronautical Journal*, 62(569), 363–376.
- Hart, I. B. (1962). *The world of Leonardo da Vinci: Man of science, engineer and dreamer of flight*. New York: Viking Press.
- Hausman, J. (2012). Contingent valuation: From dubious to hopeless. *Journal of Economic Perspectives*, 26(4), 43–56.
- Haygarth, J. (1793). *Sketch of a plan to exterminate the casual small-pox from Great Britain; And to introduce general inoculation*. London: J. Johnson.
- Health and Safety Executive. (2001). *Reducing risks, protecting people. HSE's decision-making process*. Norwich: Her Majesty's Stationery Office. <http://www.hse.gov.uk/risk/theory/r2p2.pdf>
- Heinzerling, L. (2000). The rights of statistical people. *Harvard Environmental Law Review*, 24, 189–207.
- Heinzerling, L. (2002). Markets for arsenic. *Georgetown Law Journal*, 90, 2311–2339.
- Helman, L., & Parchomovsky, G. (2011). The best available technology standard. *Columbia Law Review*, 111(6), 1194–1243.
- Hendee, W. R., & Marc Edwards, F. (1986). ALARA and an integrated approach to radiation protection. *Seminars in Nuclear Medicine*, 16, 142–150.
- Hendershot, D. C. (1997). Inherently safer chemical process design. *Journal of Loss Prevention in the Process Industries*, 10, 151–157.
- Henderson, D. A. (2011). The eradication of smallpox—An overview of the past, present, and future. *Vaccine*, 29, D7–D9.
- Hills, R. L. (1989). *Power from steam. A history of the stationary steam engine*. Cambridge: Cambridge University Press.
- Hinman, A. R. (2018). Measles and rubella eradication. *Vaccine*, 36(1), 1–3.
- Hinze, J., & Wilson, G. (2000). Moving toward a zero injury objective. *Journal of Construction Engineering and Management*, 126(5), 399–403.
- Hochkirch, A., Beninde, J., Fischer, M., Krahner, A., Lindemann, C., Matenaar, D., Rohde, K., et al. (2018). License to kill? – Disease eradication programs may not be in line with the convention on biological diversity. *Conservation Letters*, 11(1), e12370.
- Hoffman, S. (2014). Zero benefit: Estimating the effect of zero tolerance discipline policies on racial disparities in school discipline. *Educational Policy*, 28(1), 69–95.
- Holder, J., & Lee, M. (2007). *Environmental protection, law and policy. Text and materials* (2nd ed.). Cambridge: Cambridge University Press.
- Holloway, D. (2011). The vision of a world free of nuclear weapons. In C. M. A. Kelleher & J. Reppy (Eds.), *Getting to zero: The path to nuclear disarmament* (pp. 11–27). Stanford: Stanford University Press.
- Holm, H., & Sahlin, N. E. (2009). Regeringens nolivision för självmord kan få motsatt effect [The government's vision zero for suicides may have the opposite effect]. *Läkartidningen*, 106(17), 1153–1154.
- Holms, A. C. (1917). *Practical shipbuilding. A treatise on the structural design and building of modern steel vessels; The work of construction, from the making of the raw material to the equipped vessel, including subsequent up-keep and repairs*. London: Longmans, Green and Company.
- Huckman, R. S., & Raman, A. (2015). Medicine's continuous improvement imperative. *JAMA*, 313(18), 1811–1812.
- Huh, K., & Choi, C. (2016). Product recall policies and their improvement in Korea. *Management and Production Engineering Review*, 7(4), 39–47.
- Hussain, A., Ali, S., Ahmed, M., & Hussain, S. (2018). The anti-vaccination movement: A regression in modern medicine. *Cureus*, 10(7), e2919.

- Hutchins, B. L., & Harrison, A. (1911). *A history of factory legislation* (2nd ed.). London: P.S. King and Son.
- Hutchinson, D. F., & de Kerbrech, R. (2011). *RMS Titanic: 1909–12 (Olympic Class); An insight into the design, construction and operation of the most famous passenger ship of all time*. Sparkford: Haynes.
- Innes, M. (1999). ‘An iron fist in an iron glove?’ The zero tolerance policing debate. *Howard Journal of Criminal Justice*, 38(4), 397–410.
- International Commission on Radiological Protection. (1959). Recommendations of the International Commission on radiological protection: Adopted September 9, 1958. In *Annals of the ICRP* (Vol. 1) (ICRP Publication No. 1).
- International Commission on Radiological Protection. (1977). Recommendations of the ICRP: ICRP Publication No. 26. In *Annals of the ICRP* (Vol. 1, No. 3, pp. 1–53). Oxford: Pergamon.
- International Transport Forum. (2018). *Road safety annual report 2018*. Downloaded Dec. 24, 2019 from https://www.itf-oecd.org/sites/default/files/docs/irtad-road-safety-annual-report-2018_2.pdf
- Jarvis, W. R. (2007). The United States approach to strategies in the battle against healthcare-associated infections, 2006: Transitioning from benchmarking to zero tolerance and clinician accountability. *Journal of Hospital Infection*, 65, 3–9.
- Jenner, E. (1801). *The origin of vaccine inoculation*. London: D.N. Shury.
- Jensen, M. B. (2017). Radiation protection principles and development of standards for geological repository systems. In M. J. Apted & J. Ahn (Eds.), *Geological repository systems for safe disposal of spent nuclear fuels and radioactive waste* (2nd ed., pp. 601–628). Duxford: Woodhead.
- JICOSH. (n.d.). Concept of “Zero-accident total participation campaign”. <https://www.jniosh.johas.go.jp/icpro/jicosh-old/english/zero-sai/eng/index.html>. Downloaded June 23, 2019.
- Johanson, G., & Tinnerberg, H. (2019). Binding occupational exposure limits for carcinogens in the EU—good or bad? *Scandinavian Journal of Work, Environment & Health*, 45(3), 213–214.
- Johansson, R. (2009). Vision Zero – Implementing a policy for traffic safety. *Safety Science*, 47(6), 826–831.
- Johansson, M. K. V., Johanson, G., Öberg, M., & Schenk, L. (2016). Does industry take the susceptible subpopulation of asthmatic individuals into consideration when setting derived no-effect levels? *Journal of Applied Toxicology*, 36(11), 1379–1391.
- Johnston, P., & Harris, R. (2019). The Boeing 737 MAX saga: Lessons for software organizations. *Software Quality Professional*, 21(3), 4–12.
- Jones, D. D., Holling, C. S., & Peterman, R. (1975). Fail-safe vs. safe-fail catastrophes. *IIASA working paper* WP-75-93. Downloaded December 29, 2019 from <http://pure.iiasa.ac.at/id/eprint/335/1/WP-75-093.pdf>
- Kakalia, S., & Karrarb, H. H. (2016). Polio, public health, and the new pathologies of militancy in Pakistan. *Critical Public Health*, 26(4), 446–454.
- Kathren, R. L., & Ziemer, P. L. (1980). The first fifty years of radiation protection – A brief sketch. In R. L. Kathren & P. L. Ziemer (Eds.), *Health physics: A backward glance* (pp. 1–9). Elmsford: Pergamon Press.
- Kathren, R. L., Munson, L. H., & Higby, D. P. (1984). Application of risk-cost benefit techniques to ALARA and de-minimis. *Health Physics*, 47, 195.
- Kaufmann, J. R., & Feldbaum, H. (2009). Diplomacy and the polio immunization boycott in Northern Nigeria. *Health Affairs*, 28(4), 1091–1101.
- Keeley, L. H., Fontana, M., & Quick, R. (2007). Baffles and bastions: The universal features of fortifications. *Journal of Archaeological Research*, 15(1), 55–95.
- Kemp, R. (2005). Zero emission vehicle mandate in California: Misguided policy or example of enlightened leadership. In C. Sartorius & S. Zundel (Eds.), *Time strategies, innovation, and environmental policy* (pp. 169–191). Cheltenham: Edward Elgar.
- Khan, F. I., & Abbasi, S. A. (1998). Inherently safer design based on rapid risk analysis. *Journal of Loss Prevention in the Process Industries*, 11, 361–372.

- Khorasani-Zavareh, D. (2011). System versus traditional approach in road traffic injury prevention. A call for action. *Journal of Injury and Violence Research*, 3(2), 61–61.
- Kim, H. J., Haugen, S., & Utne, I. B. (2017). Reliability analysis for physically separated redundant system. In M. Ram & J. P. Davim (Eds.), *Advances in reliability and system engineering* (pp. 1–25). Cham: Springer.
- Klein, M. W. (1986). Labelling theory and delinquency policy: An experimental test. *Criminal Justice and Behaviour*, 13, 47–79.
- Kletz, T. (1978). What you don't have, can't leak. *Chemistry and Industry*, (May), 287–292.
- Kletz, T. A. (2004). Inherently safer design: The growth of an idea. *Process Safety Progress*, 15, 5–8.
- Klinke, A., Dreyer, M., Renn, O., Stirling, A., & Van Zwanenberg, P. (2006). Precautionary risk regulation in European governance. *Journal of Risk Research*, 9(4), 373–392.
- Knoll, F. (1976). Commentary on the basic philosophy and recent development of safety margins. *Canadian Journal of Civil Engineering*, 3(3), 409–416.
- Koch, T., & Denike, K. (2009). Crediting his critics' concerns: Remaking John Snow's map of Broad Street cholera, 1854. *Social Science and Medicine*, 69(8), 1246–1251.
- Kohn, L. T., Corrigan, J. M., & Donaldson, M. S. (Eds.). (2000). *To err is human, building a safer health system*. Washington, DC: Institute of Medicine, National Academy Press.
- Kretsinger, K., Strebel, P., Kezaala, R., & Goodson, J. L. (2017). Transitioning lessons learned and assets of the global polio eradication initiative to global and regional measles and rubella elimination. *The Journal of Infectious Diseases*, 216(Suppl 1), S308–S315.
- Kristianssen, A.-C., Andersson, R., Belin, M.-Å., & Nilsen, P. (2018). Swedish Vision Zero policies for safety – A comparative policy content analysis. *Safety Science*, 103, 260–269.
- Kurrer, K.-E. (2018). *The history of the theory of structures. Searching for equilibrium* (2nd ed.). Berlin: Wiley.
- Lacoe, J., & Steinberg, M. P. (2018). Rolling back zero tolerance: The effect of discipline policy reform on suspension usage and student outcomes. *Peabody Journal of Education*, 93(2), 207–227.
- Lahiri, S., Levenstein, C., Nelson, D. I., & Rosenberg, B. J. (2005). The cost effectiveness of occupational health interventions: Prevention of silicosis. *American Journal of Industrial Medicine*, 48(6), 503–514.
- Lambert, N., Strebel, P., Ornstein, W., Icenogle, J., & Poland, G. A. (2015). Rubella. *Lancet*, 385(9984), 2297–2307.
- Le Van, W. B. (1873). On the steam-boiler explosion at the rolling-mill of J. Wood & Brother, Conshohocken, PA. *Journal of the Franklin Institute*, 95(4), 248–253.
- Le Van, W. B. (1892). *Safety-valves. Their history, antecedents, invention and calculation*. New York: Norman W. Henley & Co.
- Lersow, M., & Waggitt, P. (2020). *Disposal of all forms of radioactive waste and residues. Long-term stable and safe storage in geotechnical environmental structures*. Cham: Springer.
- Lewis, E. B. (1957). Leukemia and ionizing radiation. *Science*, 125(3255), 965–972.
- Lindell, B. (1996). The history of radiation protection. *Radiation Protection Dosimetry*, 68, 83–95.
- Lindell, B., & Beninson, D. J. (1981). ALARA defines its own limit. *Health Physics*, 41, 684–685.
- Lissner, L., & Romano, D. (2011). Substitution for hazardous chemicals on an international level—The approach of the European project 'Subsport'. *New Solutions: A Journal of Environmental and Occupational Health Policy*, 21(3), 477–497.
- Lovrencic, V., & Gomišček, B. (2014). Live working as an example of electrical installation maintenance with the zero accidents philosophy. In *2014 11th International Conference on Live Maintenance (ICOLIM)* (pp. 1–8). IEEE.
- Lu, F. C. (1988). Acceptable daily intake: Inception, evolution, and application. *Regulatory Toxicology and Pharmacology*, 8(1), 45–60.
- Ma, H., Balthasar, F., Tait, N., Riera-Palou, X., & Harrison, A. (2012). A new comparison between the life cycle greenhouse gas emissions of battery electric vehicles and internal combustion vehicles. *Energy Policy*, 44, 160–173.

- Marston, D. C. (2016). Unintended consequences of 'zero tolerance' policies. <https://www.apa.org/pi/ses/resources/indicator/2016/01/unintended-consequences>. Downloaded June 26, 2019.
- Martinez, S. (2009). A system gone berserk: How are zero-tolerance policies really affecting schools? *Preventing School Failure: Alternative Education for Children and Youth*, 53(3), 153–158.
- Mashaw, J. L., & Harfst, D. L. (1987). Regulation and legal culture: The case of motor vehicle safety. *Yale Journal on Regulation*, 4, 257–316.
- McElhattan, K. (2019). Technology can help eliminate workplace deaths for good. Downloaded December 3, 2019 from <https://www.nsc.org/safety-first-blog/technology-can-help-eliminate-workplace-deaths-for-good>
- Melchers, R. E. (2001). On the ALARP approach to risk management. *Reliability Engineering and System Safety*, 71, 201–208.
- Mendoza, A. E., Wybourn, C. A., Mendoza, M. A., Cruz, M. J., Juillard, C. J., & Dicker, R. A. (2017). The worldwide approach to Vision Zero: Implementing road safety strategies to eliminate traffic-related fatalities. *Current Trauma Reports*, 3(2), 104–110.
- Merkouris, P. (2012). Sustainable development and best available techniques in international and European law. In K. E. Makuch & R. Pereira (Eds.), *Environmental and energy law* (1st ed., pp. 37–60). Chichester: Blackwell.
- Michaels, R. K., Makary, M. A., Dahab, Y., Frassica, F. J., Heitmiller, E., Rowen, L. C., Crotreau, R., Brem, H., & Pronovost, P. J. (2007). Achieving the national quality forum's 'never events': Prevention of wrong site, wrong procedure, and wrong patient operations. *Annals of Surgery*, 245(4), 526–532.
- Mielke, D. P. (2012). Fortifications and fortification strategies of mega-cities in the ancient near east. In *Proceedings of the 7th international congress of the archaeology of the ancient Near East*, Vol. 1, pp. 73–91.
- Mokkenstorm, J. K., Kerkhof, A. J. F. M., Smit, J. H., & Beekman, A. T. F. (2018). Is it rational to pursue zero suicides among patients in health care? *Suicide and Life-threatening Behavior*, 48(6), 745–754.
- Möller, N., Hansson, S. O., Holmberg, J.-E., & Rollenhagen, C. (Eds.). (2018). *Handbook of safety principles*. Hoboken: Wiley.
- Moses, F. (1997). Problems and prospects of reliability-based optimization. *Engineering Structures*, 19(4), 293–301.
- Nagaich, U., & Sadhna, D. (2015). Drug recall: An incubus for pharmaceutical companies and most serious drug recall of history. *International Journal of Pharmaceutical Investigation*, 5(1), 13–19.
- Neff, R. A., & Goldman, L. R. (2005). Regulatory parallels to Daubert: Stakeholder influence, 'sound science,' and the delayed adoption of health-protective standards. *American Journal of Public Health*, 95(S1), S81–S91.
- Nelson, E. J. (1996). Remarkable zero-injury safety performance. *Professional Safety*, 41(1), 22–25.
- Newburn, T., & Jones, T. (2007). Symbolizing crime control: Reflections on zero tolerance. *Theoretical Criminology*, 11(2), 221–243.
- Niven, C. M., Mathews, B., Harrison, J. E., & Vallmuur, K. (2020). Hazardous children's products on the Australian and US market 2011–2017: An empirical analysis of child-related product safety recalls. *Injury Prevention*, 26, 344–350.
- Nordström, M., Hansson, S. O., & Hugosson, M. B. (2019). Let me save you some time... On valuing travelers' time in urban transportation. *Essays in Philosophy*, 20(2), 206.
- Nuclear Energy Agency. (2002). *Improving versus maintaining nuclear safety*. Paris: OECD. <https://www.oecd-nea.org/nsd/reports/nea3672-improving.pdf>
- Oestreich, A. E. (2014). ALARA 1912: 'As low a dose as possible' a century ago. *Radiographics*, 34, 1457–1460.
- Ogle, R. A., Dee, S. J., & Cox, B. L. (2015). Resolving inherently safer design conflicts with decision analysis and multi-attribute utility theory. *Process Safety and Environmental Protection*, 97, 61–69.

- Ong, E. K., & Glantz, S. A. (2001). Constructing 'sound science' and 'good epidemiology': Tobacco, lawyers, and public relations firms. *American Journal of Public Health*, 91(11), 1749–1757.
- Orecchini, F. (2007). A 'measurable' definition of sustainable development based on closed cycles of resources and its application to energy systems. *Sustainability Science*, 2(2), 245–252.
- Otway, H. J., & von Winterfeldt, D. (1982). Beyond acceptable risk: On the social acceptability of technologies. *Policy Sciences*, 14, 247–256.
- Overton, T., & King, G. M. (2006). Inherently safer technology: An evolutionary approach. *Process Safety Progress*, 25, 116–119.
- Papin, D. (1681). *A new digester or engine for softening bones*. London: Henry Bonwicke.
- Parker, H. M. (1980). Health physics, instrumentation and radiation protection. *Health Physics*, 38, 957–996.
- Partington, C. F. (1822). *An historical and descriptive account of the steam engine, comprising a general view of the various modes of employing elastic vapour as a prime mover in mechanics*. London: J. Taylor.
- Pearce, D. W., Russell, S., & Griffiths, R. F. (1981). Risk assessment: Use and misuse. *Proceedings of the Royal Society of London Series A, Mathematical and Physical Sciences*, 376(1764), 181–192.
- Peierls, R. E. (1967). Zero defects by James F. Halpin [Review]. *Scientific American*, 216(1), 140–142.
- Perry, R. T., & Halsey, N. A. (2004). The clinical significance of measles: A review. *Journal of Infectious Diseases*, 189(Suppl 1), S4–S16.
- Petrosino, A., Turpin-Petrosino, C., & Guckenburg, S. (2014). The impact of juvenile system processing on delinquency. In D. P. Farrington & J. Murray (Eds.), *Labeling theory, empirical tests* (pp. 113–147). New Brunswick: Transaction Publishers.
- Polio Global Eradication Initiative. (2019). *Annual report 2018*. Geneva: World Health Organization. Downloaded December 6, 2019 from <http://polioeradication.org/wp-content/uploads/2016/07/Annual-report-2018.pdf>
- Procter, D., Young, M., & Howlett, V. (1990). Training for total quality performance: The experience of British Steel. *Total Quality Management*, 1(3), 319–326.
- Pyhälä, M., Brusendorff, A. C., & Paulomäki, H. (2010). The precautionary principle. In M. Fitzmaurice, D. M. Ong, & P. Merkouris (Eds.), *Research handbook on international environmental law* (pp. 203–226). Cheltenham: Edward Elgar.
- Randall, F. A. (1976). The safety factor of structures in history. *Professional Safety*, January, 12–18.
- Ranken, A. (1982). EPA v. National Crushed Stone Association. *Ecology Law Quarterly*, 10, 161–189.
- Rankine, W. J. M. (1859). *A manual of the steam engine and other prime movers*. London: Richard Griffin.
- Rappaport, S. M. (1993). Threshold limit values, permissible exposure limits, and feasibility: The basis for exposure limits in the United States. *American Journal of Industrial Medicine*, 23, 683–694.
- Reiman, T., & Norros, L. (2002). Regulatory culture-A case study in Finland. In *Proceedings of the IEEE 7th conference on human factors and power plants* (Vol. 5, pp. 5–20). IEEE.
- Rodgers, M. D., & Blanchard, R. E. (1993). *Accident proneness: A research review*. Civil Aero-medical Institute, Federal Aviation Administration. Downloaded December 2015 from <https://apps.dtic.mil/dtic/tr/fulltext/u2/a266032.pdf>
- Rudén, C., & Hansson, S. O. (2008). Evidence based toxicology – 'Sound science' in new guise. *International Journal of Occupational and Environmental Health*, 14, 299–306.
- Rudén, C., & Hansson, S. O. (2010). REACH is but the first step – How far will it take us? Six further steps to improve the European chemicals legislation. *Environmental Health Perspectives*, 118(1), 6–10.
- Rupp, N. G. (2004). The attributes of a costly recall: Evidence from the automotive industry. *Review of Industrial Organization*, 25(1), 21–44.

- Sabin, F. R. (1952). Trends in public health. *American Journal of Public Health and the Nation's Health*, 42(10), 1267–1271.
- Sagoff, M. (1988). Some problems with environmental economics. *Environmental Ethics*, 10, 55–74.
- Samet, J. M., & Burke, T. A. (2001). Turning science into junk: The tobacco industry and passive smoking. *American Journal of Public Health*, 91(11), 1742–1744.
- Schoenberger, H. (2011). Lignite coke moving bed adsorber for cement plants – BAT or beyond BAT? *Journal of Cleaner Production*, 19, 1057–1065.
- Shang, X., & Tonsor, G. T. (2017). Food safety recall effects across meat products and regions. *Food Policy*, 69, 145–153.
- Sherratt, F. (2014). Exploring ‘Zero Target’ safety programmes in the UK construction industry. *Construction Management and Economics*, 32(7–8), 737–748.
- Singh, J., & Singh, H. (2009). Kaizen philosophy: A review of literature. *IUP Journal of Operations Management*, 8(2), 51–73.
- Skiba, R. J. (2014). The failure of zero tolerance. *Reclaiming Children and Youth*, 22(4), 27–33.
- Skiba, R., & Peterson, R. (1999). The dark side of zero tolerance: Can punishment lead to safe schools? *The Phi Delta Kappan*, 80(5), 372–382.
- Smith, M. J., Bouch, J., Bradstreet, S., Lakey, T., Nightingale, A., & O'Connor, R. C. (2015). Health services, suicide, and self-harm: Patient distress and system anxiety. *The Lancet Psychiatry*, 2(3), 275–280.
- Snow, S. J. (2002). Commentary: Sutherland, snow and water: The transmission of cholera in the nineteenth century. *International Journal of Epidemiology*, 31, 908–911.
- Soffritti, M., Sass, J. B., Castleman, B., & Gee, D. (2013). Vinyl chloride: A saga of secrecy. In D. Gee et al. (Eds.), *Late lessons from early warnings II* (European Environmental Agency, EEA report 1/2013) (pp. 179–202). Luxembourg: Publication Office of the European Union.
- Sørensen, F., & Styhr Petersen, H. J. (1991). A process-based method for substitution of hazardous chemicals and its application to metal degreasing. *Hazardous Waste and Hazardous Materials*, 8(1), 69–84.
- Srinivasan, R., & Natarajan, S. (2012). Developments in inherent safety: A review of the progress during 2001–2011 and opportunities ahead. *Process Safety and Environmental Protection*, 90(5), 389–403.
- Stahl, S. D. (2016). The evolution of zero-tolerance policies. *CrissCross*, 4(1), 7.
- Sterman, J. D., Siegel, L., & Rooney-Varga, J. N. (2018). Does replacing coal with wood lower CO2 emissions? Dynamic lifecycle analysis of wood bioenergy. *Environmental Research Letters*, 13(1), 015007.
- Stover, J., Hallett, T. B., Wu, Z., Warren, M., Gopalappa, C., Pretorius, C., Ghys, P. D., Montaner, J., Schwartländer, B., & New Prevention Technology Study Group. (2014). How can we get close to zero? The potential contribution of biomedical prevention and the investment framework towards an effective response to HIV. *PLoS One*, 9(11), e111956.
- Stuart, R. (1829). *Historical and descriptive anecdotes of steam-engines and of their inventors and improvers*. London: Wightman and Cramp.
- Sunstein, C. R. (1991). Administrative substance. *Duke Law Journal*, 607–646.
- Svensk författningssamling. (1949). Kungl. Maj:ts kungörelse med föreskrifter angående tillämpning av arbetarskyddslagen (arbetarskyddskungörelsen) [Royal statutory regulation concerning the application of the health and safety law (the health and safety ordinance)]. Svensk författningssamling [Swedish Code of Statutes] 1949:208. Stockholm: Norstedt & Söner.
- Swuste, P., van Gulijk, C., & Zwaard, W. (2010). Safety metaphors and theories, a review of the occupational safety literature of the US, UK and the Netherlands, till the first part of the 20th century. *Safety Science*, 48(8), 1000–1018.
- Swuste, P., Van Gulijk, C., Zwaard, W., & Oostendorp, Y. (2014). Occupational safety theories, models and metaphors in the three decades since World War II, in the United States, Britain and the Netherlands: A literature review. *Safety Science*, 62, 16–27.

- Szyszczak, E. (1992). 1992 and the working environment. *Journal of Social Welfare & Family Law*, 14(1), 3–16.
- Tapping, T. (1855). *The factory acts*. London: Shaw and Sons.
- Tayeh, A., Cairncross, S., & Cox, F. E. G. (2017). Guinea worm: From Robert Leiper to eradication. *Parasitology*, 144(12), 1643–1648.
- Tengs, T. O., Adams, M. E., Pliskin, J. S., Safran, D. G., Siegel, J. E., Weinstein, M. C., & Graham, J. D. (1995). Five-hundred life-saving interventions and their cost-effectiveness. *Risk Analysis*, 15(3), 369–390.
- Thomas, P. E. (1830). \$ 4,000 premium. *Mechanics Magazine and Journal of Public Internal Improvement*, 1(12), 372–373.
- Tickner, J., Jacobs, M. M., & Mack, N. B. (2019). Alternatives assessment and informed substitution: A global landscape assessment of drivers, methods, policies and needs. *Sustainable Chemistry and Pharmacy*, 13, 1–13.
- Tingvall, C., & Haworth, N. (1999). Vision Zero – An ethical approach to safety and mobility. In *6th ITE international conference road safety & traffic enforcement: Beyond 2000*, 6–7 September 1999, Melbourne.
- Tredgold, T. (1827). *The steam engine*. London: J. Taylor.
- Trevithick, F. (1872). *Life of Richard Trevithick, with an account of his inventions* (Vol. 2). London: E. & F.N. Spon.
- Tripp, H. A. (1938). Progress in road safety measures. *Police Journal*, 11(4), 428–449.
- Tuominen, P., Reda, F., Dawoud, W., Elboshy, B., Elshafei, G., & Negm, A. (2015). Economic appraisal of energy efficiency in buildings using cost-effectiveness assessment. *Procedia Economics and Finance*, 21, 422–430.
- Twaalfhoven, S. F. M., & Kortleven, W. J. (2016). The corporate quest for zero accidents: A case study into the response to safety transgressions in the industrial sector. *Safety Science*, 86, 57–68.
- UK Chemicals Stakeholder Forum. (2010). A guide to substitution. An information note from the UK Chemicals Stakeholder Forum. Final version – 20th August 2010. https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/227031/Forum-guide-substitution-101105.pdf
- United Nations. (1992). Report of the United Nations Conference on Environment and Development, Rio de Janeiro, 3–14 June 1992, <http://www.un.org/documents/ga/conf151/aconf15126-lannex1.htm>
- United Nations, General Assembly. (1946). Establishment of a Commission to deal with the problems raised by the discovery of atomic energy. [https://undocs.org/en/A/RES/1\(I\)](https://undocs.org/en/A/RES/1(I)).
- Vägverket [The Swedish Road Administration]. (2004). Nollvisionen förändrar Sveriges vägmiljöer [Vision Zero changes the road traffic in Sweden].
- Vainio, H., & Tomatis, L. (1985). Exposure to carcinogens: Scientific and regulatory aspects. *Annals of the American Conference of Governmental Industrial Hygienists*, 12, 135–143.
- Vandenbergh, M. P. (1996). An alternative to ready, fire, aim: A new framework to link environmental targets in environmental law. *Kentucky Law Journal*, 85, 803–918.
- Vilkomir, S., & Kharchenko, V. (2012). A diversity model for multi-version safety-critical I&C systems. In *Proceedings of the international conference ESREL-PSAM11*, Helsinki, June, 2012. Downloaded December 30, 2019 from <https://pdfs.semanticscholar.org/8c4f/fb40b87eca10aacc46f131437e48a18f9aee.pdf>
- Wahlström, Bo. (1999). The precautionary approach to chemicals management: A Swedish perspective. In C. Raffensperger and J. Tickner (Eds.), *Protecting public health and the environment. Implementing the precautionary principle* (pp. 51–69). Washington, DC: Island Press.
- Wang, L. K., & Yang, J. Y. (1975). Total waste recycle system for water purification plant using alum as primary coagulant. *Resource Recovery and Conservation*, 1(1), 67–84.
- Warye, K., & Granato, J. (2009). Target: Zero hospital-acquired infections. *Healthcare Financial Management*, 63(1), 86–92.

- Warye, K. L., & Murphy, D. M. (2008). Targeting zero health care-associated infections. *American Journal of Infection Control*, 36(10), 683–684.
- Wasserman, D., Tacic, I., & Bec, C. (2022). Vision Zero in suicide prevention and suicide preventive methods, this volume.
- Watson, R., Kassem, M., & Li, J. (2019). A framework for product recall in the construction industry. In *Advances in ICT in design, construction and management in architecture, engineering, construction and operations (AECO)* (pp. 755–765). Newcastle upon Tyne: Northumbria University.
- Wein, S. (2013). Exploring the virtues (and vices) of zero tolerance arguments. In D. Mohammed, and M. Lewiński (Eds.), *Virtues of argumentation. Proceedings of the 10th international conference of the Ontario Society for the Study of Argumentation (OSSA)*, 22–26 May 2013, pp. 1–10. Windsor: OSSA. Downloaded December 15, 2019, from <https://scholar.uwindsor.ca/cgi/viewcontent.cgi?article=2141&context=ossaarchive>
- Weinberg, A. M. (1972). Science and trans-science. *Minerva*, 10, 209–222.
- Wiener, J. B., & Rogers, M. D. (2002). Comparing precaution in the United States and Europe. *Journal of Risk Research*, 5, 317–349.
- Wilson, R. (2002). Precautionary principles and risk analysis. *Technology and Society Magazine, IEEE*, 21(4), 40–44.
- Wilson, J. Q., & Kelling, G. L. (1982). Broken windows: The police and neighborhood safety. *Atlantic Monthly*, 249(3), 29–38.
- Worth, L. J., & McLaws, M.-L. (2012). Is it possible to achieve a target of zero central line associated bloodstream infections? *Current Opinion in Infectious Diseases*, 25(6), 650–657.
- Wright, E. (2012). Olympic health and safety: Record breakers. *Building*, 1 June. Downloaded December 2, 2019 from <https://www.building.co.uk/olympic-health-and-safety-record-breakers/5036956.article>
- Young, S. (2014). From zero to hero. A case study of industrial injury reduction: New Zealand Aluminium Smelters Limited. *Safety Science*, 64, 99–108.
- Zaman, A. U. (2014). Measuring waste management performance using the ‘Zero Waste Index’: The case of Adelaide, Australia. *Journal of Cleaner Production*, 66, 407–419.
- Zaman, A. U. (2015). A comprehensive review of the development of zero waste management: Lessons learned and guidelines. *Journal of Cleaner Production*, 91, 12–25.
- Zaman, A. U., & Lehmann, S. (2011). Challenges and opportunities in transforming a city into a ‘zero waste city’. *Challenges*, 2(4), 73–93.
- Zou, P. X. W. (2010). Fostering a strong construction safety culture. *Leadership and Management in Engineering*, 11(1), 11–22.
- Zwetsloot, G. I. J. M., Aaltonen, M., Wybo, J.-L., Saari, J., Kines, P., & Op De Beeck, R. (2013). The case for research into the zero accident vision. *Safety Science*, 58, 41–48.
- Zwetsloot, G. I. J. M., Kines, P., Wybo, J.-L., Ruotsala, R., Drupsteen, L., & Bezemer, R. A. (2017a). Zero Accident Vision based strategies in organisations: Innovative perspectives. *Safety Science*, 91, 260–268.
- Zwetsloot, G. I. J. M., Kines, P., Ruotsala, R., Drupsteen, L., Merivirta, M.-L., & Bezemer, R. A. (2017b). The importance of commitment, communication, culture and learning for the implementation of the Zero Accident Vision in 27 companies in Europe. *Safety Science*, 96, 22–32.
- Zwetsloot, G., Leka, S., & Kines, P. (2017c). Vision zero: From accident prevention to the promotion of health, safety and well-being at work. *Policy and Practice in Health and Safety*, 15(2), 88–100.
- Zwier, J., Blok, V., Lemmens, P., & Geerts, R.-J. (2015). The ideal of a zero-waste humanity: Philosophical reflections on the demand for a bio-based economy. *Journal of Agricultural and Environmental Ethics*, 28(2), 353–374.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.



Arguments Against Vision Zero: A Literature Review

3

Henok Girma Abebe, Sven Ove Hansson, and
Karin Edvardsson Björnberg

Contents

Introduction	108
Vision Zero: What It Is	111
Vision Zero as a Goal	111
Vision Zero as a Strategy	112
Vision Zero as New Responsibilities	113
Four Central Assumptions of Vision Zero	115
Moral Criticism	117
“It Is Morally Misguided to Strive for a World Free from Suffering”	117
“It Is Not Ethically Unjustified That People Die on the Roads”	118
“Safety Should Not Have Higher Priority than Everything Else”	120
“It Is Immoral to Focus Only on Fatal and Serious Injuries”	122
“Vision Zero Is Paternalistic”	122
“Vision Zero Goes Contrary to Equity and Social Justice”	125
Rationality-Based Criticism	129
“Vision Zero Is Unrealistic”	130
“Vision Zero Is Too Imprecise”	132
“Vision Zero Is Counterproductive and Self-Defeating”	134
Operational Criticism	136
“Accident Statistics Do Not Provide a Reliable Picture of the Safety Level”	137
“Vision Zero Neglects the Probability of Accidents”	139

H. G. Abebe (✉) · K. Edvardsson Björnberg
Division of Philosophy, KTH Royal Institute of Technology, Stockholm, Sweden
e-mail: hgirma@kth.se; karin.bjornberg@abe.kth.se

S. O. Hansson
Division of Philosophy, Royal Institute of Technology (KTH), Stockholm, Sweden
Department of Learning, Informatics, Management and Ethics, Karolinska Institutet,
Stockholm, Sweden
e-mail: soh@kth.se; sven-ove.hansson@abe.kth.se

“Too Little Responsibility Is Assigned to Drivers”	140
“Too Little Responsibility Is Assigned to System Designers”	141
Conclusion	144
Cross-References	145
References	145

Abstract

Despite Vision Zero’s moral appeal and its expansion throughout the world, it has been criticized on different grounds. This chapter is based on an extensive literature search for criticism of Vision Zero, using the bibliographic databases Philosopher’s Index, Web of Science, Science Direct, Scopus, Google Scholar, PubMed, and Phil Papers, and by following the references in the collected documents. Even if the primary emphasis was on Vision Zero in road traffic, our search also included documents criticizing Vision Zero policies in other safety areas, such as public health, the construction and mining industries, and workplaces in general. Based on the findings, we identify and systematically characterize and classify the major arguments that have been put forward against Vision Zero. The most important arguments against Vision Zero can be divided into three major categories: moral arguments, arguments concerning the (goal-setting) rationality of Vision Zero, and arguments aimed at the practical implementation of the goals. We also assess the arguments. Of the 13 identified main arguments, 6 were found to be useful for a constructive discussion on safety improvements.

Keywords

Vision Zero · Nollvisionen · Criticism · Road safety · Ethics · Systems thinking

Introduction

The adoption of Vision Zero (“Nollvisionen”) in Sweden in 1997 represented a crucial shift in road safety management (Government Bill 1996/97:137). Road safety work at the time was heavily influenced by utilitarian cost-benefit analysis and by an approach that considered failing road users to be the main cause of road accidents. In contrast, Vision Zero emphasized the responsibility of system designers and clearly prioritized safety over mobility and cost containment. It declared that the fatalities and serious injuries that result from preventable crashes are morally unacceptable. Moreover, it assumed that road users want health and self-preservation and that this is what the design and operation of the road system has to deliver. The moral appeal and relative success of Vision Zero has led to its acceptance in more and more countries, states, and cities around the world, and it has had a considerable impact also in other areas of public safety than road traffic (Mendoza et al. 2017; Kristianssen et al. 2018).

However, the global proliferation of Vision Zero policies does not imply that it is without flaws. In fact, Vision Zero has sustained a fair amount of criticism, both in academic literature and in the public debate. So far, these criticisms have not been

investigated systematically. Therefore, in this chapter we aim to identify, categorize, and critically assess the arguments that have been put forward against Vision Zero. Our categorization of arguments is based on a desk-based review of academic research articles, reports, and policy documents from the last two decades. The documents were retrieved through searches in the bibliographic databases, Philosopher's Index, Web of Science, ScienceDirect, Scopus, Google Scholar, PubMed, and Phil Papers, and by following the references in the collected documents. Even if the primary emphasis was on Vision Zero in road traffic, our search also included documents criticizing Vision Zero policies in other safety areas, such as public health, the construction and mining industries, and workplaces in general.

Our analysis shows that the most important arguments against Vision Zero can be divided into three major categories: moral arguments, arguments concerning the (goal-setting) rationality of Vision Zero, and arguments aimed at the practical implementation of the goals (see Fig. 1).

Firstly, critics target the central moral assumptions behind Vision Zero, such as its uncompromising prioritization of safety and its assumption that deaths and serious injuries in the road traffic system are morally unacceptable. For instance, the ethical assumption behind Vision Zero has been criticized by authors who claim that it is morally acceptable that some people die on the road, since driving is a risky activity that they chose voluntarily to engage in. Moreover, it has been argued that the resources required to realize Vision Zero will have to be taken from other policy areas where they could be used to greater advantage from an ethical point of view. Vision Zero has also been accused of being paternalistic and unjust, and some of the measures proposed to realize it have been accused of threatening the freedom, autonomy, and privacy of road users.

Secondly, critics question the rationality of setting and working toward the goal to prevent all fatalities and serious injuries in traffic safety. It has been argued that such a goal is unrealistic and therefore irrational to pursue. Doing so is counterproductive, according to the critics, since the agents who are responsible for achieving it will become demotivated when they realize that no matter how great effort they invest, the goal will never be achieved. In addition, Vision Zero has been criticized for being too imprecise to be serviceable as a goal for public policy.

Thirdly, criticisms target specific operationalizations of Vision Zero that have been used in its practical application. The ways in which safety is measured in the application of Vision Zero to road system design has been criticized. Some critics have claimed that too little responsibility is assigned to system designers. Others maintain that system designers are assigned too much responsibility and that this will reduce drivers' sense of responsibility and make them drive more dangerously.

In section "[Vision Zero: What It Is](#)," we introduce Vision Zero and its central assumptions. Sections "[Moral Criticism](#)," "[Rationality-Based Criticism](#)," and "[Operational Criticism](#)" present and analyze the arguments that we have found in each of the three categories just mentioned. Section "[Conclusion](#)" summarizes our findings and identifies some arguments against Vision Zero that are, in our view, particularly worthy of further consideration and analysis.

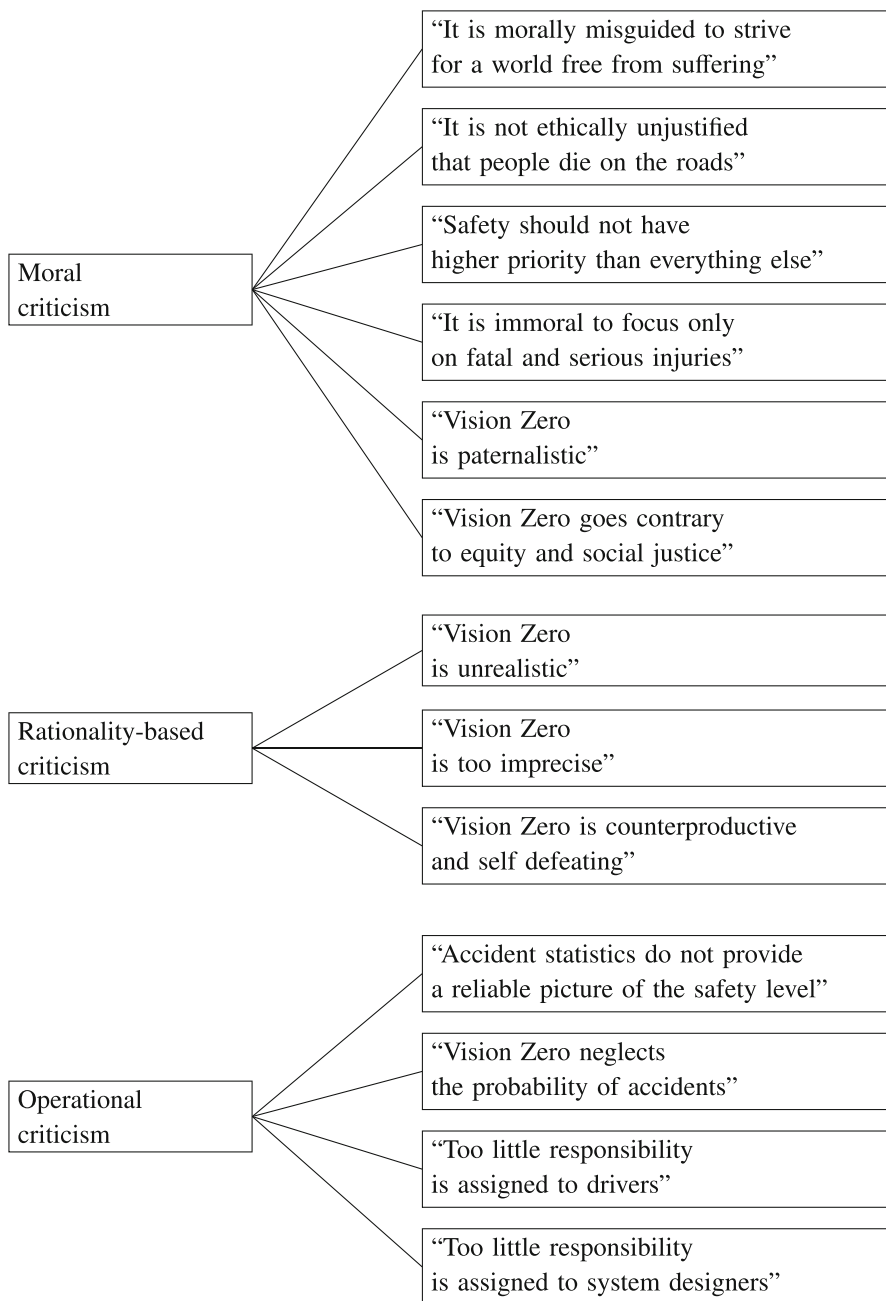


Fig. 1 The arguments against Vision Zero discussed in this chapter

Vision Zero: What It Is

A significant number of countries have adopted and are committed to Vision Zero. It was first adopted in 1997 when the Swedish parliament unanimously endorsed it as the country's traffic safety policy (Belin and Tillgren 2012). Currently, similar Vision Zero policies are in force in a number of other countries, including Finland, Norway, Denmark, the Netherlands, Germany, Poland, the UK (London), Australia, New Zealand, and Canada (see Part 2 of this handbook). While New York was the first city to adopt the policy in the USA (in 2014), many other cities have joined the group since then (Mendoza et al. 2017). So, what is Vision Zero and how does it differ from the safety policies it came to replace?

Vision Zero as a Goal

According to the Swedish government, the long-term goal of road safety is that “no one should be killed or seriously injured as a result of traffic accidents in the road transport system” (Government Offices of Sweden 2016, p. 6). Despite the government's use of the term “vision,” it is clear from the preparatory work that Vision Zero is in fact a policy goal that is supposed to guide all road safety work in Sweden (Government Bill 1996/97:137). To reach the goal, which is not temporally specified, substantial adjustments of the road transport system will have to be made over an extended period of time.

As a policy goal, Vision Zero functions not only as a symbolic expression of the government's ambition to reduce the number of fatalities and serious injuries in the road system. The goal also guides and induces action toward achievement of the desired end-state. Using terminology from goal-setting literature, the goal is “achievement-inducing” (Edvardsson and Hansson 2005). As with most policy goals, Vision Zero coordinates action both temporally and between individuals and organizations. Vision Zero can be used by the national transport administration as a departure point for developing and implementing a series of safety measures over time in such a way that the desired end-state can more easily be reached. It can also be used to allocate resources among various sub-agencies or departments to the same effect. Based on Vision Zero, implemented road safety measures can be evaluated and adjusted, and responsibility for insufficient goal achievement can be established. Thus, Vision Zero functions as a normative framework against which road safety measures can be developed, implemented, evaluated, and adjusted (Rosencrantz et al. 2007; see also ► Chap. 1, “Vision Zero and Other Road Safety Targets”, by Edvardsson Björnberg, in this handbook). In this effort, Vision Zero posits the fallibility of human beings as a starting point for the design and operation of roads and vehicles (Johansson 2009). But, importantly, Vision Zero is not only a goal but also a strategy.

Vision Zero as a Strategy

Vision Zero is a strategy that relies on both social and technological innovations in the process of approaching the goal of zero fatalities and serious injuries (Belin et al. 2012). Vision Zero differs fundamentally from the traditional approach to road safety management in terms of its “problem formulation, its view on responsibility, its requirement for the safety of road users, and the ultimate objective of road safety work” (Belin et al. 2012, p. 171).

Problem formulation and ultimate objective: In the traditional approach to road safety, traffic accidents were presented as the major problem to be solved, and individual road users were believed to be causally responsible for up to 95% of those accidents (Evans 1996). In contrast, Vision Zero puts focus not on the accidents per se but on the resulting fatalities and serious injuries. The difference between the traditional approach and Vision Zero can be clearly seen from the measures advocated by proponents of the two approaches. In Vision Zero, a road safety measure that leads to an overall decline in fatalities and serious injuries is preferable, even if it involves a greater number of accidents or minor injuries. This is, for instance, the main logic behind the shift from traffic lights to roundabouts in four-way intersections in most Vision Zero-committed countries, such as Sweden and the Netherlands (Mendoza et al. 2017). While roundabouts, as compared to traffic lights, tend to lead to a greater number of crashes, the reduced speed in roundabouts makes the crashes less severe, and the number of fatalities and severe injuries is considerably lower (ibid.). When it comes to road and street design, Vision Zero goes contrary to the traditionally dominant safety strategy of increasing space for vehicles through the construction of wider roads, wider lanes, straighter roads, and larger crossings (Bergh et al. 2003; Johansson 2009). Although these measures facilitate the flow of traffic and reduce the number of crashes, they often have negative effects on safety since “the most predominant effect of creating more space is an increase in driving speed, which means higher levels of kinetic energy in crashes” (Johansson 2009, p. 828).

Two prominent improvements in vehicle technology that have brought huge safety gains in Swedish roads are the introduction of seat belt reminders (SBR) and alcohol interlocks. A study by Krafft et al. (2006) of the driving behavior of 3000 Swedish drivers showed that “in cars without SBR, 82.3 percent of the drivers used the seat belt, while in cars with SBR, the seat belt use was 98.9 percent” (p. 125). Furthermore, “in cars with mild reminders, the use was 93.0 percent” (p. 125). From this, the authors concluded that installing seat belt reminders in all cars would have a dramatic impact on the number of fatal and seriously injured car occupants. Seat belt reminders are a prime example of a measure that aims at reducing the consequences rather than the probability of crashes.

Alcohol interlocks provide another important example of a technological innovation with huge safety benefits. Drunk driving is one of the major factors involved in crashes leading to fatalities and serious injuries. According to the WHO’s global status report (WHO 2018), between 5% and 35% of all road fatalities are alcohol-related. In Sweden and many other European countries, alcohol interlocks have been

introduced as a remedy to the problem of drunk driving. The technology is now widely employed in professional settings. In 2017, 97% of the busses operating in public transport in Sweden had an alcohol interlock (Sveriges Bussföretag 2018). The technology requires a driver to exhale into the machine and prevents the driver from starting the vehicles if a certain amount of alcohol is detected in their breath. Alcohol interlocks is one of many measures in traffic safety that have positive impacts both on the probability and the severity of crashes. Drunk drivers are more often involved in crashes, and these crashes also tend to lead to more serious injuries.

Vision Zero as New Responsibilities

In the traditional approach to traffic safety, the individual road user was identified as the most important causal factor in traffic accidents. Based on accident investigations, it was reported that road users' behavior was the cause of about 95% of traffic crashes (Evans 1996). Consequently, it was assumed that road users carry almost the whole responsibility for traffic safety, and it was often concluded that safety propaganda, rather than technical improvement, was the best way to deal with the problem.

However, these reports were based on a questionable approach to causality, and the conclusions were largely unhelpful in attempts to improve road safety. Although we usually prefer to think in terms of "the cause" of an accident or other event, the assumption of a single cause is in many cases a gross oversimplification. Events do not typically follow from one single cause. Instead, there are several causal factors, all of which contribute to the effect. Various practical considerations influence which causal factor we tend to call "the cause," for instance, how certain we are of its influence, its conspicuity, whether it could plausibly have been absent, and whether it could have been changed by human action (Hoover 1990). For instance, if you ask a bacteriologist what is the cause of cholera you can expect the answer "the bacterium *Vibrio cholerae*," but a public health expert will probably give an answer referring to the lack of proper sanitation. These causal descriptions are useful for different purposes. In the treatment of cholera patients, the answer mentioning the microorganism may be the most adequate one, whereas the answer referring to sanitary conditions is more useful for disease prevention.

In much the same way, most traffic accidents have causal factors pertaining both to the behavior of the driver and to the construction of the vehicle and the road system. For instance, a driver's decision to drive drunk is often a causal factor contributing to an accident. However, there are also various other causal factors, including the social conditions that led the driver to drinking too much, the lack of resources for treatment of alcoholism, and vehicle-related causal factors such as the lack of an alcohol interlock on the car in question. In discussions on how to reduce traffic accidents involving drunk drivers, the drivers' decisions were previously almost exclusively at focus, whereas the decisions by regulators and manufacturers to allow respectively market cars without alcohol interlocks have not been part of

the discussion. The situation was similar for other types of traffic accidents. (On causality and responsibility in road traffic, see also ► [Chap. 5, “Responsibility in Road Traffic”](#), by Hansson.)

One of the basic insights behind Vision Zero is that it is often inefficient to focus on the causal factors that have traditionally been called “the cause” of various accidents. Instead, the focus should be on the causal factors that are most accessible to interventions that improve safety. It then becomes clear that technological factors such as the construction of vehicles and roads are usually much easier to change than human behavior. This has led to a whole range of new technological solutions that have reduced the number of serious road accidents. Where individual road users fail to act or behave as they are expected to, due to factors such as negligence, incompetence, lack of knowledge, or health issues, the road system can be redesigned so that people do not die or get seriously injured even when mistakes are made. As noted by Johansson (2009, p. 827): “It is true, that 95% of all crashes or collisions depend on human error, but according to Vision Zero philosophy, 95% of the solutions are in changing roads, streets or vehicles.”

In consequence, Vision Zero has led to a new focus on the responsibilities of the governmental, regional, and local authorities that are involved in the design of the road environment, as well as the responsibilities of vehicle manufacturers. These two groups are called the system designers, and according to Vision Zero they shared the ultimate responsibility for traffic safety (McAndrews 2013; Government Offices of Sweden 2016). According to Tingvall (1997, p. 41), the road system designers “bear the responsibility to do everything in their power to make the system as safe as possible. . . they are also responsible for meeting the road user demands for road safety in the system.”

In part this is an institutional responsibility, carried by the agencies and companies that construct roads and vehicles. However, it also has an important component of professional responsibilities. The engineers and other professionals who perform the actual construction tasks have responsibilities, both individually and collectively, to make the choices that save lives and avert suffering. A comparison can be drawn with healthcare. Governments are responsible for organizing healthcare systems that save lives and preserve health. This is an institutional responsibility. At the same time, physicians, nurses, and other healthcare professionals have a responsibility – again, both individually and collectively – to make the choices that best serve the health and well-being of their patients.

The professional responsibilities in Vision Zero go beyond traditional blame responsibility (often called backward-looking responsibility), which assigns blame for causing a traffic safety problem. The main focus is put on task responsibility, which is concerned with who can do something about the problem. In Vision Zero, the overarching task responsibility falls on the system designers. But unavoidably, blame responsibility can also become involved. System designers can be held responsible for inactivity or misdirected activity that leads to fatalities or serious injuries that could otherwise have been prevented. (On responsibility ascriptions, see also ► [Chap. 5, “Responsibility in Road Traffic”](#), by Hansson.)

Responsibility is not a zero sum game. In other words, if one group takes on more responsibilities, then this does not mean that some other group has to become less responsible. The fact that system designers assume new responsibilities does not relieve individual road users of their responsibility to drive safely and respect traffic regulations (Tingvall 1997). On the contrary, in Vision Zero, the moral responsibility of road users goes beyond what was traditionally expected of them. In addition to the duty of respecting and abiding by the traffic rules and regulations, the “moral responsibility of road users extends to the health of all road users in all situations—even those not anticipated or defined by the legislative and governing bodies. The moral responsibility of road users also involves making clearly stated and powerful demands on the designers of the system” (Tingvall 1997, p. 42).

Four Central Assumptions of Vision Zero

The above discussion suggests that Vision Zero builds on a set of important but controversial assumptions, all of which are necessary to justify the adoption and promotion of the policy.

Ethical Assumption: “It Can Never Be Ethically Acceptable That People Are Killed or Seriously Injured When Moving Within the Road Transport System”

Vision Zero is based on the ethical assumption that it is morally unacceptable that people get killed or seriously injured due to preventable traffic crashes. For the proponents of Vision Zero, any goal other than zero amounts to voluntarily permitting that people are killed or seriously injured on the road (Tingvall and Haworth 1999). This ethical basis of Vision Zero is the major justification for the adoption of the policy in many Vision Zero-committed countries and cities. Importantly, it has called established practices in safety work and transport decision-making into doubt. For instance, this applies to the use of cost-benefit analysis in road safety planning, since CBA often trades the safety of road users to promote other values. Moreover, monetary valuation of human life and the use of willingness to pay in determining the economic value of traffic safety measures are deemed morally problematic from a Vision Zero perspective (Hokstad and Vatn 2008).

From this point of view, the level of road fatalities and serious injuries is the product of our choices as a political society regarding which values we should prioritize. Fatalities and serious injuries are not deemed to be necessary costs. Instead, they show what price a society is willing to pay for mobility. This is a radical departure from the traditional approach to traffic safety, in which traffic fatalities and injuries are viewed as the necessary costs of using the road system (Belin et al. 2012). Unlike the traditional approach to traffic safety in which safety is usually compromised to promote mobility, Vision Zero considers such a compromise as an unsatisfactory situation that should be changed. According to Tingvall (1997, p. 56):

It goes without saying that human life cannot be exchanged for some gain. To give an example, if a new road, new car design, new rule etc. is judged as having the potential to save human life, then the opportunity must always be taken, provided that no other more cost-effective action would produce the same safety benefit.

Empirical Assumption: Human Fallibility Is Unavoidable and Therefore Has to be Taken into Account in Traffic Safety Work

There is a long history from industrial safety of attempts to avoid accidents by identifying the workers who cause them and taking measures aiming at these individuals. However, this strategy has been found to be inefficient, since accidents are not limited to the actions of a special category of particularly accident-prone individuals. Therefore, industrial safety instead focuses on making operations “fail-safe,” or “inherently safe,” which means essentially that the prevalence of human mistakes is accepted and focus is put on minimizing the negative consequences following from such mistakes (Hansson 2010; Hammer 1980; Harpur 1958; Jones et al. 1975). A similar development has taken place in patient safety, where a “blame culture” looking for scapegoats has largely been replaced by a focus on how the probabilities and the consequences of such mistakes can be reduced (Rall et al. 2001).

Vision Zero can be seen as representing the same trend, applied primarily to traffic safety. Traditionally, the mistaken behavior of individual road users was taken to be the dominant cause of safety problems in the road traffic system. Consequently, traffic rules and regulations, education, training, licensing, and other mechanisms for behavioral change were emphasized, with the pronounced intention of promoting the required behavior and adjusting the road user to the road system (Belin et al. 2012). Vision Zero instead focuses on making the road system “fail-safe,” so that human mistakes do not lead to serious accidents. This approach is based on the simple observation that, in contrast to human nature, vehicle technology and road infrastructure are accessible to radical change.

Operational Assumption: The Ultimate Responsibility for Traffic Safety Should be Assigned to System Designers

This assumption has largely the same motivation as the previous one. From a Vision Zero perspective, the ultimate cause of accidents is taken to be the “imperfect system.” Therefore, it is the system that needs to be adjusted to the needs and capabilities of the individual road users, not the other way around. Since the problem of traffic safety is systemic in nature (Larsson et al. 2010), Vision Zero presumes that responsibility should be shared among the actors that directly or indirectly influence the safety of the system.

Empirical Assumption: Technology Can Solve Most Road Traffic Safety Problems

In most countries that have shown a significant improvement in traffic safety over the past few decades, the role of technology has been significant. The introduction of seat belts, seat belt reminders, airbags, automatic brakes, alcohol interlocks,

motorcycle and bicycle helmets, and safer road and street designs have played and continue to play a key role in preventing fatalities and injuries. It is generally believed that further progress can be achieved with new, innovative technologies. However, the use and application of most of the technologies that improve traffic safety has long been questioned and debated due to their impact on economy, freedom, autonomy, and privacy. Nonetheless, in countries committed to Vision Zero, a strong emphasis on the development and implementation of innovative technologies appears to be the next step. The Swedish Vision Zero recommends the use of the best available technology when addressing road safety problems, hence emphasizing the role of technological innovation in promoting traffic safety. In the USA, one of the three major strategies identified in *The Road to Zero: A Vision for Achieving Zero Roadway Deaths by 2050* (Ecola et al. 2018) is to accelerate the production and use of advanced technologies.

Moral Criticism

We will consider six moral arguments against Vision Zero. Four arguments claim that Vision Zero assigns too high priority to serious injuries in road traffic. These arguments are presented in order of decreasing strength of the claims that they make. We discuss the argument that Vision Zero is paternalistic and in section “[Vision Zero Goes Contrary to Equity and Social Justice](#)” the argument that it counteracts social justice.

“It Is Morally Misguided to Strive for a World Free from Suffering”

It has been argued that, because Vision Zero aims to achieve zero fatalities and serious injuries through the categorical prioritization of safety and health of road users, it seeks to create a risk-free society, which is considered problematic in various ways. Firstly, there is the argument that creating a risk-free society conflicts with individual liberty, interpreted as the freedom of individuals to choose what risks they wish to expose themselves to (see section “[Too Little Responsibility Is Assigned to Drivers](#)”). Ekelund (1999), for instance, criticized Vision Zero for aiming to eliminate all road traffic risks despite the fact that some people are willing to take more risks than others. In the context of public health policy, Fugeli (2006) similarly argued that Vision Zero is a luxurious quest of rich European countries to create a risk-free, perfect society. In his view, Vision Zero seeks to purify life and remove defects and risks, which will lead to undesirable consequences. What these authors seem to argue is that by adopting and pursuing Vision Zero policies society may well reduce suffering in the form of deaths and serious injuries caused by certain activities, such as driving, but it also denies people the opportunities of enjoying life to a fuller extent than what is possible under a Vision Zero regime.

Dekker et al. (2016) locate Vision Zero within the “Western Judeo-Christian salvation narrative,” i.e., “the notion that a world without suffering is not only

desirable but achievable, and that efforts expended toward the goal are morally right and inherently laudable” (p. 219). This narrative understands human suffering as the result of bad choices made by individuals. Consequently, suffering can be relieved by hard work and better individual choices. This is in line with much traditional safety work, according to which the causal responsibility for accidents is largely attributed to the individual. However, Dekker et al. argue, aiming to relieve suffering by focusing on individual choices invites gaming – both by individuals, who in employment settings may refrain from reporting injuries for fear of being blamed, and managers and CEOs, who may refrain from reporting incidents that may lead to the loss of bonuses – and creates more suffering in the end.

The claim that Vision Zero seeks to achieve a perfect society is not backed up by any evidence. We have found no indication of any such assumption in the written documentation on Vision Zero. On the contrary, a major assumption behind Vision Zero is the recognition that traditional approaches to traffic safety, criticized by Dekker et al. (2016), have failed in their relentless attempts to create a perfect road user. (Cf assumption 2 in section “[Four Central Assumptions of Vision Zero](#)”) Vision Zero differs from this approach in accepting the occurrence of mistakes, and hence even accidents, as an inevitable fact of life. This speaks strongly against the claim that Vision Zero aims to create a perfect society, free from any suffering. It is difficult to imagine a totally risk-free society, constituted of imperfect individuals who are by their own nature liable to make mistakes and act on the basis of wrong judgments. Furthermore, Vision Zero does not aim at eradicating all accidents and injuries but only those that will lead to “an unacceptable loss of health” (Tingvall and Haworth 1999). Non-serious traffic injuries are outside the scope of Vision Zero. Therefore, as was rightly indicated by Zwetsloot et al. (2013, 2017), the criticism that Vision Zero seeks to create a risk-free society is more of a misconception than a genuine argument against it.

In summary, the argument that Vision Zero errs in trying to create a perfect society is based on a blatantly incorrect description of Vision Zero, and not worth taking seriously. (Therefore, we do not see a need to discuss another assumption underlying this argument, viz., that attempts to move in the direction of a “perfect” state are condemnable.)

“It Is Not Ethically Unjustified That People Die on the Roads”

One of the underlying assumptions behind the adoption and promotion of Vision Zero policies is the claim that it is morally unacceptable that people die and get seriously injured due to predictable and preventable crashes. Therefore, Vision Zero is “presented as a more, or perhaps the only, ethically sound approach” (Elvebakk 2005, p. 18). However, Elvebakk argues, the mere ambition to prevent all fatalities and serious injuries cannot in itself justify the ethical superiority of Vision Zero because “there are not necessarily major differences between wanting to reduce the number of serious accidents as much as possible, and wanting to eradicate them altogether. It would seem that either way, the best one can do is one’s best”

(Elvebakk 2005, p. 21). Moreover, “it is not necessarily *in itself* ethically unjustifiable to allow hundreds of people die in traffic every year. [...] Death is, after all, a fact of life, and as a society we have to accept that people will die, for one reason or another” (Elvebakk 2005, pp. 24–25).

Elvebakk goes on to present examples of cases of fatalities and serious injuries in different aspects of human life, where the causalities, she argues, are often deemed morally acceptable because of the mere fact that those who died or were injured had voluntarily and knowingly chosen to engage in activities associated with considerable risk. Examples are deaths as a result of suicide, drug overuse, skiing, fishing, swimming, etc. Although these risky activities claim a significant number of lives every year, Elvebakk claims that “there are relatively few calls for regulation, as risk seems to be accepted as an integral part of the activity” (Elvebakk 2005, p. 25). For her, these different areas or activities, including road traffic, belong in the “private space,” where individuals often voluntarily and knowingly choose to engage in risky activities and accept responsibility for doing so. Elvebakk comments:

Proponents of vision zero prefer not to compare road traffic to these areas, but to other professional fields, where fatalities are typically not deemed acceptable. But, arguably, the road traffic system cannot be straightforwardly compared to these professional areas, as they belong to different spaces: road traffic is (for most of the drivers) not a professional space. (Elvebakk 2005, p. 25)

Allsop (2005) advances a quite similar view regarding the nature of the road system and road users’ responsibility. For him too, the road system is not a “closed system in which everything can be defined as someone’s contractual responsibility, but as part of everyone’s day-to-day lives, which they expect to be largely free to lead” (p. 15). Moreover, Allsop identifies an additional similarity between these other risky activities that people often engage in and road traffic: most of them serve the same purpose of fulfilling and giving meaning to human life. Most people who lose their lives due to involvement in one of these risky activities have engaged in it “to meet either social needs, or demands for goods, or desires for fullness of life” (ibid.). Using the roads, he says, serves similar purposes. He concludes that “neither in terms of rational socioeconomic policy nor in terms of human desire for fulfillment is it unacceptable in principle for use of the roads to involve some risk of death or serious injury” (ibid.).

These arguments do not take into account that most of those who are killed and seriously injured in road traffic did not wish to take any risks. They just had no other choice than to travel in the risky traffic system that we have. Furthermore, the assumption that a risk is unproblematic if it comes with a voluntarily chosen activity is quite problematic. On the face of it, humans may seem to choose risk-taking. However, people taking risks do not usually desire the risk per se, but rather something else that it is associated with. For instance, consider a person who chooses to bungee jump. Arguably, what she is looking for is not the risk of dying or being seriously injured, but rather an advantage that it is associated with, namely, the thrill,

not the risk. If she had the choice of an otherwise exactly similar jump but with a safer cord, then she would presumably choose the safer alternative (Hansson 2013). The same seems to be the case for dangerous behavior in road traffic, such as speeding and drunk driving. These activities are undertaken for various reasons, including the pursuit of thrill (in the case of speeding). The claim that people drive dangerously because of a wish to increase the probability that they will end up in a wheel chair or a coffin is not borne out by any empirical evidence or plausible argument. To this should of course be added the observation that most dangerous behaviors in road traffic impose risk on other road users. We therefore have good reasons to write off the argument that we might as well let people die on the roads since they have taken the risks themselves.

“Safety Should Not Have Higher Priority than Everything Else”

The adoption of Vision Zero was partly a reaction to the use of cost-benefit analysis (CBA) in transport policy and decision-making. (See Hokstad and Vatn (2008) and Hansson (2007) for elaborate discussions on the moral and philosophical issues associated with use of CBA.) Unlike CBA, Vision Zero does not promote the weighing of safety against other values in the traffic system. Life and health, it is claimed, “can never be exchanged for other benefits within the society” (Tingvall and Haworth 1999, p. 2).

Proponents of Vision Zero have claimed that it rectified a previous double standard for different transport systems. Safety had the highest priority in aviation and rail traffic, where accidents were treated as unacceptable events. In contrast, accidents in the road system were taken to be unavoidable and a price worth paying for mobility (Johnston et al. 2014). The high demands on airplane safety have seldom been criticized, and no attempts seem to have been made to systematically evaluate safety measures in that area with cost-benefit analysis. In contrast, the application of a similarly strict attitude to road traffic, which is promoted as part of Vision Zero, has attracted much criticism. Elvik (1999) maintained that the uncompromising prioritization of safety and health in the road traffic system would divert economic resources from other societal objectives to the promotion of road safety. Since resources are limited, he argued, this would reduce measures against other causes of death and injury in society, leading to an increase in general mortality. For similar reasons, Elvebakk maintained that from a utilitarian point of view, “rather than being a more ethical approach to road safety, vision zero is a less ethically sound basis for policy” (Elvebakk 2005, p. 24). Allsop (2005) argued that “the cold socio-economic logic of the human mind and the warm aspiration of the human spirit join their voices to say: no, they are not paramount, and yes, they can be traded. [...] Safety is for living; living is much more than just keeping safe” (p. 15).

Nihlén Fahlquist (2009) argued that Vision Zero could potentially be used to justify radical limitations of freedom of movement and individual autonomy and that it could lead to privacy infringements if inbuilt technologies and safety/surveillance cameras store data on drivers’ behavior.

This criticism is based on the assumption that Vision Zero implies that traffic safety always has a higher priority than everything else. That is a misunderstanding. Proponents of Vision Zero accept that it cannot immediately be fully implemented. If traffic safety had higher priority than everything else, then all road traffic would have to be stopped immediately and only be restarted to the extent that it could be undertaken with no risk of fatalities. However, contrary to proponents of CBA, defenders of Vision Zero do not treat trade-offs, for instance, between safety and economy as optimal and satisfactory states. Instead, they treat such trade-offs as temporary compromises that should as soon as possible be superseded by new arrangements ensuring improved safety.

This can be clarified by a comparison with other social goals. There are a large number of policy areas in which society has goals that are subject to compromises with other goals. However, the relationship between goal-setting and compromises is different for different areas. In some areas, the tradition is to work with goals that are believed to be fully attainable. Economic policies illustrate this practice. It would be highly desirable to eradicate unemployment, but economic and labor market policies are not conducted in terms of such goals. Instead, more realistic goals are used, in this case a reduction in unemployment that is considered to be compatible with other goals for economic policies. In other areas, goals are used that represent the most desirable state rather than a compromise. For instance, law enforcement policies do not aim at an economically optimal frequency of manslaughter. Instead, they are based on the assumption that every case of manslaughter is one too much. Similarly, agencies for workplace health and safety are not instructed to try to achieve an economically optimal frequency of fatal accidents on workplaces but to reduce their number as much as possible. The difference between these two approaches is shown in Fig. 2. Either we make compromises and adjustments first and then set the goals (as in economic policies) or we set goals first and make compromises afterward (as in law enforcement and workplace safety). Vision Zero can be seen as an attempt to transfer traffic safety from the first to second of these patterns. This does not mean that the avoidance of traffic fatalities will be the only social goal that is never subject to trade-offs. Instead, it means that Vision Zero will be one of several goals that are given so high priority that any trade-offs will be treated as temporary and unsatisfactory concessions.

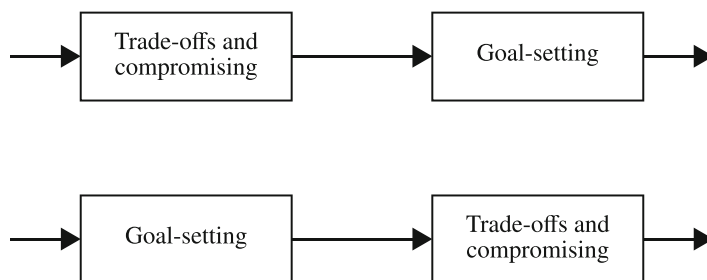


Fig. 2 Two approaches to goal-setting and compromising

In this perspective, the argument that Vision Zero crowds out all other social goals is essentially a straw man argument. However, since the relationship of Vision Zero to other social goals is seldom sufficiently clarified, this is a criticism that has the virtue of giving rise to useful and clarifying discussions.

“It Is Immoral to Focus Only on Fatal and Serious Injuries”

One important point where Vision Zero differs fundamentally from traditional safety approaches is its problem formulation (Belin et al. 2012). As noted above, the traditional goal of road safety was to prevent accidents, regardless of how severe they were. In contrast, Vision Zero accepts that accidents are inevitable in a complex system filled with cognitively fallible individuals. Therefore, it is argued, the road system should be forgiving, and so constructed that predictable crashes do not have severe consequences. Notably, crashes are often not a result of conscious negligence of instituted traffic rules and regulations but of honest and minor errors of judgment (Elvebakk 2007). Another reason for emphasizing fatalities and serious injuries in road safety is, of course, that it is those accidents that bear the largest personal, social, and economic costs.

In a recent book criticizing the Vision Zero approach in Victoria, Australia, Morgan (2018) identifies some debatable aspects of Vision Zero’s emphasis on fatal and serious injuries. Singling out and focusing on such crashes, he argues, fails to take into account the magnitude of suffering caused by minor injuries and the economic cost associated with them. He claims that “fatal and serious injury crashes are only a small part of the total road safety/vehicle collision problem” (Morgan 2018, p. 48).

It is fairly easy for a defender of Vision Zero to address this argument. It is generally accepted that saving lives has a much higher priority than preventing accidents that will only lead to temporary impairments of health and mobility. Furthermore, it can be argued that the focus on severe accidents was a crucial factor for making Vision Zero realistic enough to be adopted as a national traffic safety policy in several countries. However, it should also be conceded that the avoidance of minor accidents cannot be given zero priority. Although there does not seem to be a need to give up the strong priority for avoiding fatalities and serious injuries, there is certainly a need to discuss how less serious accidents can be included in preventive work that has a Vision Zero framework as its major driving force.

“Vision Zero Is Paternalistic”

There is a long history of criticism against safety measures in road traffic that are perceived as restricting individual liberty. Legislation against drunk driving has been a major target of such criticism and so have seat belts and bicycle and motorcycle helmets (Jones and Bayer 2007; McKenna 2007). One major argument that is usually presented against the promotion of such safety interventions is that they

tend to diminish individual autonomy and pose undue interference in an individual's personal life. Much of this criticism has been couched in anti-paternalist terms (► Chap. 6, "Liberty, Paternalism, and Road Safety", by Hansson). It has been argued that as long as no harm is done to others, individuals should be allowed to do what they voluntarily choose to do in road traffic. This type of criticism has repeatedly been directed at Vision Zero. Ekelund (1999) argues that people who so wish should be allowed not to use safety belts, helmets, or other safety technologies. Allsop (2005) maintains that Vision Zero is morally problematic due to the restrictions it imposes on individuals seeking to engage in activities that make their life complete and meaningful, even at the expense of losing their health and safety. Elvebakk (2015) has presented what is probably the most extensive criticism along these lines. She provided two major reasons why road traffic systems operating in accordance with Vision Zero will be problematic from the viewpoint of individual liberty.

The first reason is related to the responsibility ascriptions in Vision Zero. Traditionally, individual road users almost exclusively took the blame for accidents. Moreover, the road system was conceived as a private sphere of individual road users, where they could act and behave as they wanted, so long as they took responsibility for their actions and behavior (Elvebakk 2007). According to Vision Zero, however, it is the responsibility of the system designers to design a road system that takes into account the fallibility and physical vulnerability of road users. Individual road users will still be responsible for respecting traffic rules, but "if they do not live up to these expectations, the system designers must take measures" (Nihlén Fahlquist 2009, p. 391). This, Elvebakk claims, means that contrary to previous systems in which road users themselves could determine the level of risk they wanted to take, in Vision Zero only the system designers determine the level of risk in the system. This argument is obviously fallacious since it is based on the incorrect assumption that road users in a traditional system can choose the level of risk they are exposed to. Many of the people who have been killed on the roads drove as carefully and safely as they could but were hit by another vehicle that suddenly appeared in a place where it should not be. This applies not least to pedestrian and cyclist fatalities.

Elvebakk's second argument is based on the observation that if the intention in Vision Zero is to bring down the number of killed and injured to zero, then system designers cannot allow road users to engage in "high risk activities" in the road traffic system. This observation is correct, and it is also true that proponents of Vision Zero have proposed and partially implemented measures that restrict the liberty to engage in high-risk activities on the road, such as speeding and drunk driving. The use of alcohol interlocks, seat belt locks, and intelligent speed adaptation (ISA) will have a significant impact on the safety of the road system. According to Elvebakk (2015, p. 301):

Although these technologies only reinforce existing regulation, they do in fact represent a considerable reduction of the individual road user's actual autonomy: while a ban merely adds a legal risk to the existing risk of the action, a coercive technology – if successful –

physically prevents the individual from carrying out the undesired action. Thus, to the extent that the measures are introduced to protect the road users performing the undesired action, they do take paternalism to a significantly higher level.

When evaluating this criticism, it is important to note that few if any of the measures proposed to implement Vision Zero are in fact paternalistic. For instance, Elvebakk commits a serious conceptual mistake when claiming that the introduction of alcohol interlocks is an expression of paternalism. According to the Global Road Safety Partnership (2007), the presence of even small amounts of alcohol in drivers' blood increases the risk of being involved in crashes. A recent report by the International Transport Forum shows that more than 273,000 annual deaths in the road traffic systems are alcohol-related (Vissers 2017). Obviously, a drunk driver poses a risk not only to her- or himself but also to other users of the road system. For instance, a report by the Centers for Disease Control and Prevention (1997, p. 104) indicates that "approximately one fourth of all traffic deaths among children aged <15 years involved alcohol and that in nearly two thirds of passenger deaths involving a legally drunk driver, the child was in the car driven by the legally drunk driver."

Alcohol interlocks, as well as speed limits that are also essential components of Vision Zero implementation, restrict the freedom of drivers to drive as they wish. However, the issue at play is not:

My freedom to drive as I like
versus
Public measures to protect me.

Instead it is:

My freedom to drive as I like
versus
Public measures to protect others on the roads and pavements.

Thus, criticism against Vision Zero for being paternalistic is largely misdirected. It is not paternalistic to prevent a person from engaging in an activity that exposes others to risks of death. It should be noted that even before Vision Zero, major reductions in the number of road traffic casualties had been achieved with non-paternalistic measures that restrict individual liberty. This includes requirements of licenses, speed limit laws, and drunk driving laws. Technological measures that further reduce the prevalence of speeding and drunk driving, such as alcohol interlocks and automatic speed adapters, certainly infringe on the liberty to behave in certain ways on the roads, but these measures are by no means paternalistic. It may be rhetorically efficient to defend the liberty to put others' lives in danger by labeling countermeasures as paternalistic, but this is certainly not a valid argument.

According to McKenna (2007), an important lesson from the experience with such interventions is that the perceived legitimacy of an activity and the associated intervention determine both the implementation and final success of the

intervention. McKenna uses the example of how difficult it was to succeed with interventions against drunk driving in the past, when it was perceived to be a morally acceptable practice, albeit illegal. However, as the public perception of drunk driving shifted from acceptance to considering it to be an antisocial activity, the preconditions for implementing interventions also changed; it became easier for law enforcement bodies to take “active steps to detect and deter drunk driving” (McKenna 2007, p. 2). As this shows, the perceived legitimacy of an activity can change over time. What is considered legitimate at one point in time may not remain so over time. In a study performed in Sweden, Norway, and Denmark, Eriksson and Bjørnskau (2012) investigated the public’s acceptance of three ICT-based traffic safety measures that have implications on the privacy and freedom of individual road users. The measures were speed cameras, intelligent speed adaptation (ISA), and event data recorder (EDR). The study indicated that awareness of the problem for which the intervention is used, the belief that one’s own actions could contribute to addressing the safety problem, belief in the fairness and effectiveness of the measure, and demographic factors influenced the acceptance of these measures. Generally, the study reported relatively high levels of acceptance for all three measures, despite their impact on privacy and freedom for the drivers concerned.

In summary, the argument that Vision Zero is paternalistic does not get off the ground, since the major restrictions on drivers’ behavior that have been proposed to implement Vision Zero are all non-paternalistic. (On paternalism and traffic safety, see also ► Chap. 6, “Liberty, Paternalism, and Road Safety”, by Hansson.)

“Vision Zero Goes Contrary to Equity and Social Justice”

Globally, the burden of road traffic fatalities and injuries is disproportionately borne by pedestrians, bicyclists, and motorcyclists, who account for more than half of all deaths on the road. It has now been established that road traffic injury is the leading cause of death for children and young adults aged 5–29 years. According to the WHO, a major reason for this is that road safety planning and decision-making usually ignore the interests of these groups (WHO 2018). In many parts of the world, vulnerable road users are forced to use the same roads as vehicles operating at speeds that can lead to fatality or a serious injury if a crash occurs. In addition to the inequitable distribution of risks between different groups of road users, the measures taken to address the problem of road safety can impact differently on different segments of a population. Safety interventions tend to be instituted mainly in areas where people can afford them, which means that investments in safety tend to favor the rich (Elvik 2003). Moreover, when road safety policies are implemented in areas distinguished by large socioeconomic gaps, there is a risk that the policies, rather than addressing the road safety issue equitably, will further exacerbate the unequal state of affairs. While such concerns are almost nonexistent in, for example, a Swedish context, much has been written about traffic-related inequity in the USA, mainly in New York City (NYC).

The most serious of these criticisms are directed against the continued use of intensive policing as a safety intervention in the Vision Zero regime. Lee (2018) argues that Vision Zero has become an essential part of systematic segregation and discrimination in the streets of NYC. In his view, Vision Zero has been repurposed to serve a system of white supremacy that relies heavily on the policing of people of color to create a safe space for rich white people. These observations are made in relation to what he calls Vision Zero apartheid. Much of his criticism is directed toward the New York Police Department (NYPD) and the way they approach electric bike (e-bike) riders. Despite not causing many injuries, Lee argues, the City and NYPD have been using Vision Zero to police and ticket mostly immigrant delivery workers. To take an example, in 2017 over 923 e-bikes were confiscated from immigrant delivery workers and nearly 1800 e-bike criminal court summonses were issued, according to Lee (2018). Criminal court summons is particularly troublesome for immigrant workers, Lee notes, since if they do not show up in court, an arrest warrant will be issued for them.

Vision Zero, as initially developed in Sweden, clearly prioritized the prevention of fatalities and serious injuries and hence excluded minor injuries and noninjury crashes from consideration. The major justification was that it is impossible to avoid all crashes, given the fundamental fact that road users are cognitively fallible. The actual reality on the ground is very different, according to critics of Vision Zero in NYC. The police still target and penalize road users who commit low-level offenses that are not interesting from a Vision Zero point of view. Moreover, in the case of delivery workers on e-bikes, they do so despite lack of credible scientific evidence linking the use of e-bikes by the delivery workers to a serious loss of health (Lee 2018). According to Lee, the targeting of the delivery workers by the police is rather designed to “calm white fears of non-white bodies by using enforcement to impose punitive forms of racial and social control under the guise of public safety” (Lee 2018, p. 186). Thus, he continues, the policing strategy is just an extension and manifestation of systemic discrimination and bias against people of color and immigrants by enforcement agencies.

The enforcement strategies of NYC and NYPD must be understood against the background of the long history of policing in the USA, where a main strategy to prevent bigger criminal offenses has been through the intensive targeting and penalization of minor offenses (Lee 2018; Conner 2016). This policing strategy, called the “broken windows approach,” or “broken taillight policing” when applied in traffic safety enforcement, emphasizes the targeting of minor offenses with the view that this prevents people from engaging in major crimes. According to Conner, the continued use of this strategy has led to a situation:

where a violation relatively insignificant to safety is aggressively and subjectively enforced. The results are the disparate stopping, ticketing and arresting of drivers and bicyclists in predominantly African-American neighborhoods. Broken taillight policing criminalizes nonviolent and non-criminal behavior, and thus risks creating opposition to enforcement against dangerous driving. Further, because the summonses and arrests that result are tried in a racist criminal justice system, investigatory traffic stops are inherently inequitable. (Conner 2016, p. 16)

Conner further claims that it is impossible to achieve the Vision Zero goal without finding a proper solution to racial biases in police enforcement work and the justice system. This, it is rightly argued, is because the presence of racial discrimination in police enforcement work will lead to the misdirection of scarce public resources, “perpetuating linked cycles of racial bias and ineffective traffic enforcement” (Conner 2016, p. 18).

Connected to the criticism of the disproportionate nature of police enforcement is the issue of procedural justice when it comes to decision-making in road safety work. Critics argue that decision-making on police priorities and strategies is performed in ways that exclude affected parties and their interests. Lugo (2015) identified four major problems that Vision Zero implementation in US cities should address in order to be successful. First, she argued that Vision Zero is a Eurocentric policy, copied from Northern Europe and implemented without taking local realities and voices in the USA into account. Second, Vision Zero’s heavy reliance on police enforcement not only fails to consider the history of police violence against people of color in the USA but also opens opportunities for the police to further apply their biases. Lugo stated:

White people may look to police as allies in making streets safer; people of color may not. . . It really doesn’t seem like Vision Zero was designed to admit the problems that are an unfortunate reality for many in this country, a reality that other groups are working very hard to bring to light. It’d be great to see the development of a street safety strategy that starts with a dialogue on what “safety” means and whose safety we have in mind, taking it for granted that we don’t all face the same safety problems. (Lugo 2015, p. 3)

The assumption that most people of color would not opt for increased policing as an intervention appears to have some empirical support. A case study on Portland City’s Vision Zero equity efforts by the Vision Zero Network shows that community stakeholders and partners who were consulted on the issue of policing were not in favor of “increased penalties and fines for traffic violations” or the use of “check-points and saturation patrols to control for DUIs,” mainly due to fear of police racial profiling (Vision Zero Network 2018, p. 3).

The third problem with the Vision Zero initiative that Lugo identified is what she calls combative issue framing. The presentation of Vision Zero as “the only ethical choice,” Lugo claims, is meant to shame politicians by suggesting that disagreeing with the vision is unethical. However, Lugo urged that this could also have detrimental “silencing effects” on already oppressed people.

I’ve seen a worrying tendency among bike advocates to dismiss those who disagree with them as NIMBYs, flattening opposition regardless of whether it comes from community members who lived through the ravages of urban renewal or privileged homeowners concerned about an influx of colored bodies into their suburban sanctum. Vision Zero strategists should show their respect for meaningful inclusion through welcoming intersectional perspectives. (Lugo 2015, p. 2)

Last but not least Lugo criticized Vision Zero proponents’ “emphasis on top-down strategy,” where the main responsibility to bring about the required change

is delegated to policy makers and planners, overshadowing the importance of initial inclusion of other stakeholders in the policy process. According to Lugo, this “creates well-known barriers to participation in agenda setting by the very users the projects. . . are intended to serve” (Lugo 2015, p. 2).

Similar concerns of exclusion of affected parties from decision-making are aired in Lee’s (2018) research on immigrant delivery workers:

Despite the sizeable presence of delivery cyclists, city officials and bike planners and advocates do not involve delivery cyclists in dialogue about street safety and design. Partly, planning processes typically privilege top-down technocratic decision-making that discounts the embodied knowledge of people and communities particularly marginalized ones. (Lee 2018, p. 46)

These criticisms concern the way decisions are made and who is involved in the decision-making processes in Vision Zero. In modern democracies, deliberation by concerned stakeholders on a proposed piece of legislation or policy action is a requirement before the legislation or intervention is put into effect. If there are parties that could be affected by the legislation or action, then involvement and consultation of these parties is an important step that determines not only the legitimacy and acceptability of the legislation or action but also its success.

Generally, when discussing the issues of equity and social justice in Vision Zero, it is important to note, as mentioned briefly earlier, that some countries and cities committed to Vision Zero inherited a road traffic system that is highly characterized by inequitable distribution of benefits and burdens in the road system. These realities have two major implications for Vision Zero when it comes to ensuring the promotion of equity in traffic safety work.

First, it is essential to identify the sources, nature, and extent of past and present inequity and to determine how they now affect the promotion of equity in Vision Zero safety work. For instance, the US General Accounting Office (GAO) in 1983 and the United Church of Christ Commission for Racial Justice in 1987 both confirmed the primary role of race and economy in the distribution of environmental benefits and burdens in the USA. Later studies have also confirmed this to be the case (Bullard 1990; Bullard and Wright 2009). In such sociopolitical environments, it is important for Vision Zero efforts to recognize the impact that race and economy could have on the distribution of benefits and burdens in the road system. Discrimination on the basis of race or economy manifests itself, for instance, through lack of recognition for people’s concerns in public decision-making and also through denying them the opportunity to meaningful participation in decision-making processes on issues that affect their lives. Hence, according to social justice scholars (Young 1990; Schlosberg 2007), the correction of distributional inequity calls for consideration and inclusion of these components of justice, which have previously been ignored but are highly important in determining who gets what in a society. Generally, these theorists claim that distributional problems could not be grasped without recognizing other important aspects that determine the processes and

outcomes of distribution. For instance, they present recognition and participation as important aspects of justice. It is argued that lack of recognition and exclusion from decision-making processes causes unfair distributive results. These considerations are particularly important in countries and cities where race and economy have a large influence on the distribution of benefits and burdens. Moreover, promoting equity in Vision Zero could also require measures to correct past injustices and unfair distributions through mechanisms such as compensation, or reforming of legal and sociopolitical institutions that could have contributed to the inequitable distribution in the first place. In the USA, for instance, we currently see a growing call for compensating previously neglected areas through increased budgets for traffic safety work. Moreover, there is a similar interest in reforming public institutions such as enforcement agencies that have long and complicated relationships with people of color, minorities, and the economically disadvantaged (Morse 2015). It is also important that Vision Zero proponents design and implement strategies for equity and make sure that current safety work does not result in unfair distribution of benefits and burdens. Conner rightly comments that:

for all cities adopting Vision Zero, an intersectional and inclusionary equity analysis must permanently guide engineering, education and enforcement along the lines of age, gender, geography and socio-economic condition as well as race. Equity must become a fourth “E,” applied in a recurring process of analysis, implementation, and evaluation. Achieving equity in Vision Zero is not only a moral obligation; equity is a tool and tactic requisite to reach our goal. (Conner 2016, p. 18)

To conclude, the criticism against Vision Zero for perpetuating inequalities is valid, although not as a criticism against Vision Zero as such but as a criticism against implementation practices, in particular in places with an entrenched history of discrimination. As we see it, this is a criticism that should be taken seriously. Countries and cities committed to Vision Zero have the double burden of addressing the causes and ill effects of past transportation injustices and making sure that decision-making and policy implementation in the Vision Zero era result in an equitable and fair outcome for all.

Rationality-Based Criticism

A second category of arguments against Vision Zero concerns the rationality (rather than the moral justification) of adopting and pursuing the goal to prevent all fatalities and serious injuries in road traffic. We discuss the argument that Vision Zero is unrealistic and, thus, cannot be used to guide and motivate action toward the desired end-state of no fatalities or serious injuries. After that we discuss the argument that Vision Zero is too imprecise to guide action effectively. Finally, we address the argument that Vision Zero, partly because it is an unrealistic and to some degree imprecise goal, is counterproductive, or self-defeating.

“Vision Zero Is Unrealistic”

A common argument against Vision Zero is that it is a utopian or entirely unrealistic goal: no matter how much we try, we will never be able to reach a state where nobody is killed or seriously injured on the roads. When the Swedish government's ministry memorandum on Vision Zero was sent out for referral in the late 1990s, a few of the consultation bodies brought up the issue of achievability. Among those critical to Vision Zero were the county council of Jämtland and Täby municipality, both of which argued that Vision Zero was unrealistic given the extensive economic and administrative resources that would be required to achieve the goal (Government Bill 1996/97:137, section “[Accident Statistics Do Not Provide a Reliable Picture of the Safety Level](#)”). A report published by the Swedish National Road and Transport Research Institute (VTI) in 2005 confirmed that similar views were held by local politicians in the mid-2000s (Roos and Nyberg 2005). In this study, in-depth interviews were conducted with 20 municipal politicians responsible for road safety work regarding their views on road safety and the implementation of Vision Zero measures. A core finding was that a majority of politicians considered Vision Zero to be important but unrealistic. However, the practical implications of holding such views were not clear-cut. A few of the interviewed politicians emphasized that it was meaningless to have a vision that was impossible to achieve. Others, however, maintained that Vision Zero was nevertheless the only morally acceptable goal to pursue.

The achievability of Vision Zero has been questioned also in the academic literature. In relation to workplace safety, Long (2012, p. 27) claimed that “absolute goals, regardless of their excuse as aspirations, break the first rule in the fundamentals of the psychology of goal setting – achievability.” In Long's view, while adoption of realistic goals typically fosters trust in the achievability of the goal and primes the agent for success, adoption of overly difficult goals leads to skepticism and instead primes the agent for failure. Similarly, in his criticism of Vision Zero traffic safety policy in the State of Victoria, Australia, Morgan (2018) argued that the goal of zero fatalities and serious injuries is impossible to achieve. Based on case studies on fatal and serious injury crashes in six areas over the period of 2012–2016, Morgan concluded that even when the widespread use of vehicle technology (autonomous braking) is realized, “some 25% to 30% of all fatal and serious crashes are still unlikely to ever go away, even with reduced urban speed limits.” However, Morgan does not cite any publications providing details of these studies. In the absence of detailed data, it is not possible to assess to what degree they support his conclusions.

In the goal-setting literature, attainability is often put forward as a rationality criterion for goals (Edvardsson Björnberg 2008). Goals need to be achievable, it is argued, in order to have the capacity to guide and motivate agents toward the desired end-state expressed by the goals. Thus, the SMART criteria, a set of goal criteria commonly referred to in management literature, include the requirement that goals should be attainable. One of the main arguments supporting this conclusion is that

goals that are utopian, or very difficult to achieve, risk becoming counterproductive. That is, when the agent realizes that she will not be able to reach the goal, her motivation to pursue it will taper off. Instead of stimulating action toward goal achievement, the goal could make it more difficult to reach or approach the desired end-state (Hansson et al. 2016). (This argument is further discussed in section “[Vision Zero Is Counterproductive and Self-Defeating](#)”)

There are at least two possible counterarguments to the “anti-utopian objection” raised by Long (2012) and others. Firstly, although empirical evidence supports the conclusion that totally unrealistic goals can have a demotivating impact (see below), a binary categorization of goals as either realistic or unrealistic is too simplistic for most policy purposes. It fails to acknowledge that goal achievability often comes in degrees. A goal that is utopian in the sense of having a very small chance of ever being fully achieved can nevertheless be approached to a meaningful degree. Many of the political goals fought for throughout history, such as equality and freedom, are in fact goals that may never be fully achieved but can still be approached to a meaningful degree. Thus, Rosencrantz et al. (2007, p. 564) write:

ideological goals like these cannot be achieved once and for all, but will always have to be fought for. This does not prevent social and political movements from using ideals such as freedom and justice as goals. It does not seem constructive to claim that goals like these should never be set, but should be replaced by goals that are known to be fully achievable. The only demand of attainability that seems to be generally required is that goals should be approachable, i.e., it should be possible to increase the degree to which they are achieved.

Highly ambitious goals are commonly adopted, not only by political decision-makers; they also play an important role in private organizations. As an example, Kerr and LePelley (2013) discussed the introduction of “stretch goals” by General Electric’s then CEO Jack Welch in the early 1980s. Inspired by Japanese-style management techniques, Welch was convinced that highly ambitious goals should be adopted in order to stimulate creativity, exploratory learning, and “outside-the-box thinking” among the company’s employees. Since then, several other companies have introduced a similar approach to goal-setting, among them the US Southwest Airlines and Toyota (Sitkin et al. 2011).

Secondly, as argued in section “[It Is Not Ethically Unjustified That People Die on the Roads](#),” there may be ethical reasons why the goal of achieving zero fatalities and serious injuries should be retained, even if it may well be impossible to fully achieve. Some political goals are difficult to adjust without losing their moral appeal. Consider, for instance, the goals to abolish slavery or human trafficking. There are good reasons for arguing that, from an ethical point of view, no number of slaves or trafficking victims above zero is good enough for these societal ambitions. In our view, the same argument applies to Vision Zero. As long as there are measures that can be taken to reduce the number of fatalities and serious injuries in road traffic, Vision Zero can be considered a rational goal.

“Vision Zero Is Too Imprecise”

Goals typically need to be precise in order to have the capacity to guide and coordinate action effectively. Vision Zero has been criticized for failing also on this account. For instance, Lind and Schmidt (2000) argued that the strategy behind Vision Zero is vague and difficult to relate to, especially for actors at regional and local levels, since it has not been operationalized into more concrete targets and measures. One suggested solution to this problem is to introduce subsidiary goals in road safety work. This has been done in some Vision Zero countries, for example, Sweden, where the overall goal of zero fatalities and serious injuries was operationalized into the more precise sub-goal to reduce the number of road traffic fatalities to 220 by 2020. (With 223 dead on Swedish roads in 2019, the country was close to achieving this sub-goal (Transport Styrelsen 2020).)

Elvebakk and Steiro (2009) investigated how the Norwegian Vision Zero was interpreted and perceived among those working with transport and road safety in the country, including politicians, representatives of the National Public Roads Administration and the Council for Road Safety and Police, and NGOs. They concluded that:

the interpretative flexibility of the vision and relative lack of public debate have created a situation where actors focus on different aspects of the vision, and on different levels, from theoretical questions of ethics to specific practical questions of implementation. On the whole, it seems that the connection between the different levels of the vision are somewhat tenuous, and in this situation actors are relatively free to construct their own interpretation, rather than building one shared vision. (Elvebakk and Steiro 2009, p. 958)

A similar attempt to explore how Vision Zero is conceptualized and instantiated by key actors in Norway was made by Langeland (2009). Among other things, this study investigated how Vision Zero policy documents address the problem of conflicting goals and interests. One of the problems of adopting nonspecific goals, identified by the author, is that responsibility for addressing potential goal conflicts is transferred from the political level (where it arguably ought to be handled) to the administrative level, where different agencies may prioritize differently in the absence of clear political directions:

By keeping the zero vision on an abstract level, the actors may evade the conflicts that will arise when it is instantiated. The actors might find this beneficiary, as it gives them more leeway. When the zero vision is instantiated, conflicting interests and competing goals come to the fore. This may generate uncertainty for the parties involved. The more the zero vision is instantiated in terms of actual change, the more difficult it will become to ensure implementation. When the zero vision is instantiated through new policies, it will challenge goals competing with road safety. This will probably impede further realization of the zero vision. (Langeland 2009, p. 76)

There can be no doubt that lack of precision can decrease the action-guiding capacity of a goal. Imprecise goals can be difficult to follow. They can also be difficult to evaluate. However, the degree of goal specificity required for a goal to

guide and coordinate action effectively depends on the context in which the goal is implemented. For instance, in a context where the implementing agents have fairly good knowledge about what to do in order to reach or approach the goal, the goal does not have to be as precise as when such knowledge is lacking. Furthermore, it is important to recognize that trade-offs may have to be made between the action-guiding and motivating properties of a goal, since a goal that has a high degree of precision may not be particularly motivating and vice versa. In practice, the action-guiding and motivating aspects of goals often have to be balanced in goal-setting processes.

In general, goals that are implemented by another actor than the goal-setter require a greater degree of precision. Edvardsson and Hansson (2005) distinguish between three different types of precision: directional, complete, and temporal precision. A goal is directionally precise if it tells the agent in which direction to go in order to approach the goal. Complete precision means that it is in addition clear to what extent the goal should be reached. A goal is temporally precise if it includes a specified point in time when it should be achieved. Directional imprecision appears to be particularly deleterious, since it leaves the agent without a clear view of what to do in order to approach the goal. In organizational contexts, where goals are adopted and implemented by actors at different levels, imprecision typically also makes it more difficult to evaluate implemented measures and hold those responsible who have impeded goal achievement.

One could argue that the Swedish Vision Zero fulfills two of the three identified aspects of precision (Rosencrantz et al. 2007). Vision Zero is directionally precise, since it clearly states that there should be a reduction in the number of killed and seriously injured people on the road. It has complete precision, since it clearly aims to achieve a total prevention of fatalities and serious injuries. At the same time, the goal lacks in temporal precision; it does not indicate a precise point in time when it is to be fully achieved. However, although Vision Zero has both directional and complete precision, the emphasis on reduction of negative outcomes as an indication of safety has been criticized.

In a study of the formalization of the Swedish system designers' responsibilities between 1997 and 2009, Belin and Tillgren (2012) argued that, although the shift in responsibility ascriptions from individual road users to system designers presented a substantial change in road safety work, the change was nevertheless ambiguous. The reason for this was that it was difficult to get a clear idea of who the system designers were and exactly which of their activities ought to be regulated. Moreover, the authors suggested that, although there was a unanimous consensus on Vision Zero when it was formulated and legally adopted, conflicts of interests emerged during the implementation phase when different actors attempted to translate the vision into concrete action. This was particularly noticeable as perceptions of the costs and benefits of legislating on system designers' responsibility became more real to the stakeholders. These observations point to a fourth type of goal precision not covered by Rosencrantz et al.' (2007) tripartite definition of goal precision, namely, precision in the division of responsibility.

In summary, the empirical evidence indicates that the criticism of imprecision in the formulation of Vision Zero is apposite and also highly constructive. It shows that an overarching goal like Vision Zero is in need of more precise sub-goals that add precision in the dimensions in which the overarching goal is not precise enough for action guidance. In the case of Vision Zero, it is important that such sub-goals specify the temporal component of precision, i.e., clarify when various task should be completed. In many cases, the division of responsibility is also in need of specification in sub-goals.

“Vision Zero Is Counterproductive and Self-Defeating”

Goals are typically adopted in order to achieve (or maintain) certain states of affairs. However, sometimes goals turn out to be self-defeating in the sense that instead of furthering the desired end-states, the goals interfere with progress, making it more difficult to achieve those end-states. As noted by Hansson et al. (2016), various mechanisms can contribute to making a goal self-defeating. We have found two major types of claims that Vision Zero goal is self-defeating, referring to economic and psychological mechanisms, respectively.

Elvik (1999) warned against economic self-defeating mechanisms of Vision Zero. Measures not subjected to cost-benefit calculations would become too expensive, and the policy would end up not only being economically counterproductive but also contributing to increased mortality.

An objective of eliminating a certain cause of death, like traffic accidents, may be so expensive to realise that it reduces resources available to control other causes of death and thus increases general mortality. (Elvik 1999, p. 265)

One of the basic assumptions underlying Elvik’s argument is that there is a causal relationship between income per capita and general mortality, particularly that there is a negative relationship between income and mortality. By disregarding CBA, Elvik argued, proponents of Vision Zero seek to invest in safety measures that do not generate returns on the invested capital, and this leads to a decline in income that would be required to prevent other causes of death in the society. Moreover, Elvik (2003) conducted an investigation into the efficiency of safety policies in Sweden and Norway and claimed to have found that the road safety policies in both countries were inefficient in improving road safety. His recommendation was that making policy priorities on the basis of CBA would lead to greater improvement of safety, than priorities based on Vision Zero.

Elvik’s economic criticism is based on a so-called risk-risk analysis, i.e., a comparison between two options, both of which are expressed in terms of risk. Some risk analysts have seen this type of comparison as a way to bypass the common psychological reluctance to value nonmonetary goods in money: “Converting all health outcomes into death-risk equivalents facilitates cost-effectiveness analysis by calculating the cost per statistical life equivalent saved, and it addresses concerns

with respect to dollar pricing” (Viscusi et al. 1991, p. 34). The most common way to perform this conversion has been to employ the correlation between health and wealth. Richer people tend to be healthier and live longer. Therefore, “the critical income loss necessary to induce one fatality” (Lutter and Morrall 1994, p. 44) has been calculated and used to translate regulatory costs into fatalities. Elvik’s analysis from 1999 is an example of this approach. However, this translation is based on highly uncertain assumptions (Hansson 2017). Since regulations also give rise to business opportunities, costs of regulation cannot be equated with losses in total income. Furthermore, the fact that people tend to live longer in richer societies depends on complex and largely unknown social mechanisms. In particular, there is a strong positive correlation between longevity and income equality. Any conversion of gross national product into gains in longevity is therefore severely misleading (Neumayer and Plümper 2016). There is no ground for assuming that an economic loss anywhere in the economy gives rise to a proportionate increase in total morbidity or mortality.

The second type of self-defeatance identified in the literature relates to the motivational, or behavioral, effects of Vision Zero. As noted above, goals are achievement-inducing not only because they guide and coordinate action toward the desired end-states. Goals can also help us achieve the desired end-states by inducing, or motivating, actions that bring us closer to the goals. This is an important aspect of goal-setting, commonly referred to in psychological and behavioral research. Significant empirical evidence supports the so-called goal-difficulty function, i.e., given certain conditions (such as that the agent has the ability to reach the goal and is committed to it), more ambitious goals will typically induce greater efforts by the agent (Locke and Latham 2002). This holds true up to a certain point where the discrepancy between the goal and the agent’s actual performance will be so great that the goal no longer has the capacity to create a corrective motivation to change the agent’s behavior toward the goal. If, at that point, the goal gives rise to frustration and resignation instead of inspiration and motivation, then the goal has become motivationally self-defeating (Hansson et al. 2016).

According to some critics, Vision Zero is a good example of a motivationally self-defeating goal. For instance, Long (2012) claimed that pursuing the goal of zero harm in the mining and construction industries has negative motivational consequences that ultimately lead to its own subversion and failure:

Unachievable goals drive frustration, cynism and negativity; that in themselves diminish effort, energy, resilience and persistence. Absolutes are not achievable with humans, only for machines and gods, and even machines decay and wear out in time. (Long 2012, pp. 24–25)

The stated reason why goals drive such frustration and negativity is that they prime people, in Long’s case employees of the mining and construction industry, for failure:

Zero harm language is not neutral and leaders should be far more aware of how such language ‘primes’ workers psychologically and culturally [...] This is the problem with

zero harm language, it's non-motivational, noninspirational and counterintuitively primes workers for failure. (Long 2012, pp. 30–31)

Fugeli (2006) similarly claimed that a public health policy based on Vision Zero thinking is problematic because it promises and demands “too much” (p. 268) and eventually leads to a distressed, dangerous, and sick society. He argued that Vision Zero’s “obsessive preoccupation with risk” will create a situation where “life becomes surrounded by dangers that the zero missionaries will rescue us from” (p. 268). According to Fugeli, “the Zero-vision demands not merely zero risk, it desires zero deviation from the ideal state of mind and body. . . . Before the Zero-vision a wise furrow, sorrow, shyness, big rump, falling penis—were regarded as natural phenomenon belonging to the mixed state of being human. In the light of the Zero-vision these occurrences become medical deviations claiming restoration by hormones, drugs and knives.” In this way, he says, the Zero Vision also contributes to the generation of injustice by dividing and ruling the society for the interest of the educated elites who have “the power to define the golden standards of human life and health and to point derisively at what we will not endure and whom we will not tolerate.” However, as far we can see, this is criticism of a straw man. We are not aware of any proponents of Vision Zero who would subscribe to this interpretation of what it means.

There is another way in which Vision Zero has been criticized for being self-defeating, namely, by creating a safety culture within the organizations responsible for implementing the goal that is not conducive to the goal’s achievement. Sherratt and Dainty (2017), for instance, argued that Vision Zero, instead of promoting safety, fosters the development of a non-learning culture in which discussions and debates about safety are eliminated. This, they argued, can lead to the “zero paradox,” i.e., by adopting and working toward Vision Zero, fatal or serious life-changing accidents actually become more likely. However, judging by the intense and mostly highly constructive debates that Vision Zero has given rise to in traffic safety organizations around the world, it is difficult to see how this could be an impending danger.

In summary, none of the proposed mechanisms that would make Vision Zero counterproductive and self-defeating has been shown to have any impact in practice. Furthermore, the success in many countries of safety work based on Vision Zero speaks against the existence of any strong such mechanisms.

Operational Criticism

We have identified four operational arguments, i.e., arguments concerning the practical methods applied in implementing Vision Zero. The first of these concerns the use of accident statistics and the second the (allegedly insufficient) use of probabilistic information. The last two arguments concern Vision Zero’s distribution

of responsibilities. According to one line of argument, more responsibility should be assigned to the road users. According to another, responsibility should instead be further shifted toward system designers.

“Accident Statistics Do Not Provide a Reliable Picture of the Safety Level”

In safety work based on Vision Zero, the degree of safety is measured and evaluated in terms in the number of fatalities and serious injuries that occur. Several authors have criticized the use of this measure (Reason 2000; Long 2012; Dekker 2017). According to Long (2012, p. 18):

Zero harm, if set as a goal is an avoidance goal. One knows goal success by the absence of something rather than the presence of something. Avoidance goals are not only not positive but are not inspirational (Moskowitz and Grant 2009). Avoidance goals tend to be punitive in nature. Performance goals are much more positive and successful. In the framework of understanding motivation and learning leaders should be talking much more in cultural discourse about ‘keeping people safe’ than ‘preventing harm’. Later discussion shows how such discourse ‘primes’ others. Why does the safety community think that avoidance goals are so inspirational?

We are not aware of any evidence or plausible argument supporting the contention that avoidance goals are not inspirational. Furthermore, it is difficult to find a goal that cannot be expressed in either way. In WW2, the resistance movements in the countries occupied by the Nazis were fighting for the “avoidance goal” not to be under occupation, which could also be described as the “positive goal” to live in a free country. Vision Zero is usually expressed as the “avoidance goal” that no road user should be killed or seriously injured on the road, but it can also be expressed as the “positive goal” that everyone travelling on the roads should travel safely. Ergo, if there is a problem with avoidance goals, then it seems to be solvable with a simple reformulation.

However, there may be more to this. According to Reason (2000, p. 4), the fact that safety is often “defined and measured more by its absence than by its presence” is a safety paradox. He argued that the standard definition of safety, freedom from risks and dangers, fails to take into account the substantial features of safety. For him, safety is better presented if it is positively defined as the ability to deal with risks and hazards so as to avoid damage or losses while still achieving one’s goals. However, even more problematic than the way safety is defined, he argued, is that safety is measured in terms of the number of accidents or incidents: “An organisation’s safety is commonly assessed by the number and severity of negative outcomes (normalised for exposure) that it experiences over a given period” (p. 5). According to Reason, this is problematic for two reasons. First, it fails to recognize that there is only a weak relationship between an organization’s “safety health” and the registered negative outcomes, as chance plays a significant role in the occurrence of accidents.

As long as hazards, defensive weaknesses and human fallibility continue to co-exist, unhappy chance can combine them in various ways to bring about a bad event. That is the essence of the term 'accident'. Even the most resistant organizations can suffer a bad accident. By the same token, even the most vulnerable systems can evade disaster, at least for a time. Chance does not take sides. It afflicts the deserving and preserves the unworthy. (Reason 2000, p. 5)

Second, he argued, a decrease in accident rates does not necessarily mean that an organization's safety culture is improving. Such a decrease can be the result of instituting mandatory safety technologies or systems that resulted in an early improvement in safety. In most organizations accident rates decline rapidly in the beginning, and "then gradually bottom out to some asymptotic value" (p. 5). Once the asymptote is reached, says Reason, "negative outcome data are a poor indication of its ability to withstand adverse events in the future" (p. 5). He claims that although the presence of high accident rates implies a bad state of safety, low asymptotic values do not necessarily show good safety even though that is how such values have usually been interpreted. Such an erroneous interpretation, he indicates, is the cause of most safety paradoxes and poses practical implications that could negatively impede the promotion of safety.

Similar criticisms have been put forward by Dekker (2017), who also discussed problems associated with defining the goal of Vision Zero in terms of its "dependent variable," i.e., reduced accident outcomes, rather than independent variables that positively or negatively affect the negative accident outcome. According to Dekker, this is one of the reasons why little is known about what activities and mechanisms underlie the reduced negative outcomes achieved by Vision Zero-committed companies. For Dekker, a reduced negative outcome could just be the result of the fraudulent manipulation of the dependent variable (accident statistics), especially if improved statistical outcomes are associated with positive incentives.

Defining a goal by its dependent variable tends to leave organizations about what to do (which variables to manipulate) to get to that goal. Workers too can become too skeptical about zero sloganeering without evidence of tangible change in local resources or practices. (Dekker 2017, p. 169)

Dekker also claimed that the emphasis on the eradication of accidents often denies the real suffering of the individual workers by inviting data manipulation, stigmatization of workers involved in accidents, and the suppression of bad news. This can result in a work environment that considers mistakes as "shameful lapses, moral failures or failures of character in practice that should aim to be perfect" (Dekker 2017, p. 243). According to Dekker and Pitzer (2016), the reason why many industries face the plateauing of safety performance and, at times, get exposed to surprise fatal accidents is to be found in the very nature of the organizational structure and practices put in place to manage safety. Based on a review of relevant safety literature, they argued that organizational structures characterized by "safety practices associated with compliance, control and quantification" (p. 7) are prone to plateauing and surprise accidents. This, they say, is because in such organizations

safety performance close to zero can lead to “a sense of invulnerability,” deflection of resources into unproductive or counterproductive initiatives, continued application of obsolete practices, and the suppression of reporting of accidents that actually occurred in the organization.

These authors are right that in general, even if deaths or serious injuries are the main targets, measuring their occurrence may not be the best way to evaluate safety. This is because safety is concerned with the risk of future accidents, which may be of a different type. This is important in industries where rare but very large accidents are the major concern, such as nuclear reactors and many chemical industries. For instance, if day-to-day workplace safety is high in a nuclear reactor – no slippery floors, safe procedures for welding, low radiation exposure, etc. – this does not prove that the risk of a nuclear meltdown is also very low. The measures needed to prevent such an accident are quite different from those needed for more mundane workplace safety issues, and their success is not guaranteed by a low frequency of workplace accidents. The nuclear industry is rather extreme in this respect, but on most workplaces there is a need to carefully analyze the possibility of rare accidents or “surprise accidents.” Arguably, this is less important in road traffic than in most other areas of safety work, due to the exceptionally high yearly toll of fatal accidents that provide ample statistical material for priority-setting. However, rare but large accidents such as the collapse of a bridge or a hillside road, or a tunnel fire, surely need to be taken into account even if they do not show up in the accident statistics. Taken as a reminder of this, the criticism referred to above is relevant and should be taken into account in applications of Vision Zero.

“Vision Zero Neglects the Probability of Accidents”

Morgan (2018) argued that Vision Zero is based on a simplistic account of risk because risk is understood solely in terms of the severity of crashes and does not take into account the likelihood that crashes will occur. He writes:

The safe system approach looks at only half the equation—it does not concern itself with likelihood. . . . The safe system premise that safety is everything . . . inevitably leads to this illogicality: mobility has no value and crash likelihood is not a consideration. . . . I think it takes a distorted view of humanity and a messianic view of one’s own understanding of life to put the safe system approach to speed management. (Morgan 2018, p. 90)

Not only is Vision Zero based on a flawed definition of what risk is, Morgan argues, it is also a system that does not trust drivers as it seeks to impose a population-wide measure on actions to be committed by one in ten people. In comparison to Vision Zero, speed design principles such as the 85th percentile would render better results since they involve a level of trust in drivers. He claims that “the only benefit of the safe system approach to speed management is that it paves the way for the whole sale proliferation of automated speed cameras, as urged by the safe system manifesto” (Morgan 2018, p. 91).

This criticism is based on the assumption that Vision Zero implementation is focused exclusively on the severity of accidents and does not take their probabilities into account. This assumption is not correct. Many of the measures promoted in Vision Zero have large effects on the probability of accidents. For instance, alcohol interlocks and speed limitations reduce the risks of all kinds of accidents, and roundabouts and central barriers reduce the risk of serious accidents. Probably the most clear examples of measures that reduce the severity of accidents without reducing their probability are seat belts and bicycle helmets, both of which were introduced long before Vision Zero.

“Too Little Responsibility Is Assigned to Drivers”

Ekelund (1999, pp. 44–45) argued that Vision Zero’s responsibility ascription is counterproductive, since it puts too great emphasis on the responsibility of system designers. This, he argues, may lead to more reckless behavior by road users. The argument is part of Ekelund’s defense of the traditional emphasis on individual responsibility of road users, which he sees as an expression of the freedom of individuals to choose their own goals in life and decide which risks are worth taking:

By passing a new law for instance about bicycle helmets, instead of leaving the decision to the individual, the responsibility of individuals for their own safety is undermined. This will in practice send a signal: ‘You do not need to find out yourself about risks and make your own decisions. We have already found out the risks and made the decisions for you.’ By extension, this can induce people to make the assumption that everything that is not forbidden is safe. It will just not be worth the trouble to keep oneself informed about risks, since the government has probably already investigated the conditions of safety. This may very well result in an increased prevalence of careless behavior. (Ekelund 1999, p. 18, authors’ translation)

Hence, according to Ekelund, if a government introduces safety legislation against certain dangers, then this will lead the public to be less cautious in relation to other risks. If this were so, then we should, for instance, expect that seat belt legislation has made people more willing to climb dangerous ladders and that the extensive legislation on aviation safety should have induced people to skate on thin ice and swim in strong currents. He provides no evidence of this effect, and we are not aware of any reason to believe that it exists.

However, there are reasons to be concerned that safety legislation can lead to less responsible and more careless behavior *in the specific context* to which the legislation in question applies. For instance, it is much more plausible that measures to increase traffic safety will make drivers feel safer and therefore behave less cautiously, than that these measures will decrease the use of safety equipment in sport activities.

Grill and Nihlén Fahlquist (2012, p. 121) asked if there were “reasons to believe that ascribing responsibility for accident prevention to system designers will in fact make drivers feel less responsible for their driving and so less cautious?” They

argued that there are indeed areas where a shared responsibility could mean less responsibility for each party, such as when a certain safety device implanted in a vehicle takes over a task that would have been performed by the driver, had the safety device been absent. They presented examples from aviation where the airplane operator's familiarity with safety devices had led to inattention and complacency (Perrow 1999, pp. 152–154). In road traffic, they argued, similar effects could result from safety devices that take over a certain task from the driver and work continuously through the whole journey, such as a collision avoidance system: "Technical systems that are very sophisticated and where almost all safety hazards are guarded by automatic systems can erode the operator's feeling of responsibility" (Grill and Nihlén Fahlquist 2012, p. 121). In their article, the authors discussed the introduction and application of alcohol interlocks as a manifestation of the responsibility of system designers and refuted the criticism that the use of interlocks will make drivers irresponsible. In their view, the use of alcohol interlocks will not diminish the responsibility of the drivers because the interlock does nothing more than establishing the sobriety of the driver; it merely establishes whether the driver is sober before she can start the engine.

This test has no direct effect on the driving experience. It does not at all guarantee that the driver is a good one or that the safety of the driver and of other road users is automatically protected. There are many other safety features and conveniences in cars that do make drivers more passive, such as automatic transmission, cruise control and automatic braking systems. The interlock, on the other hand, only prevents people above a certain degree of intoxication from driving and is itself passive during the journey. (Grill and Nihlén Fahlquist 2012, p. 122)

In conclusion, it seems reasonable to assume that some but not all measures taken to reduce the occurrence of severe injuries in road traffic can have negative effects on drivers' sense of responsibility. This is therefore a criticism that should be taken seriously, as attention to it can improve the efficiency of a Vision Zero strategy.

"Too Little Responsibility Is Assigned to System Designers"

According to Vision Zero, system designers should take the overall responsibility for designing a road system in which fatalities and serious injuries will not occur. Road users are still expected to abide by traffic safety rules and regulation. Failure to follow safety rules and standards could have legal implications. Unlike the individual road users, however, no legal responsibility for safety has been assigned to system designers so far, despite the fact that they have the overall responsibility for the safety of the road system.

Belin and Tillgren (2012) have studied attempts made in Sweden during the years 1997 to 2009 to make system designers formally responsible. Based on evidence collected from official documents, they looked into the progress of the legislative process intended to formalize the responsibility of system designers. They argued that the process of formalizing the designer's responsibility involves a long and

complicated process and that there are important factors that limited the government's attempts to realize it. Unlike the initial process that led to the adoption of Vision Zero by the Swedish Parliament, in which the different stakeholders almost unanimously supported the policy, the process of formalizing the responsibility of system designers was characterized by conflicts of interest. These conflicts resulted from the perception that the benefits and costs associated with formalizing the responsibility of system designers were not fairly distributed. This, Belin and Tillgren argued, is in turn a result of a narrow conception of system designers as involving just "the state, the municipalities, and individual road administrators" (p. 94). They argued that "in such a case, we have moved to a position where the benefits are distributed to all road users, while the costs are concentrated among road administrators" (p. 94) and hence resistance against formalizing responsibility among those who perceived that they would receive an unfair share of the burden. The study also identified other factors that prevented the realization of legal responsibility for designers. These included the difficulties associated with changing the traditional responsibility ascription for traffic safety, which is well rooted in both national and international laws, the implementation of other government efforts that had similar effects as that of regulating the responsibility of designers through law, and processes and efforts at other government levels. As an example of the latter, they indicated the positive impact that the process of regulating government agency vehicles and transport services had had on enhancing the responsibility of system designers. The regulation of road administrators' safety responsibility through an EU directive also meant that Swedish road system designers were legally responsible for at least parts of the road network, i.e., the trans-European road network that passes through Sweden. In conclusion, based on the abovementioned reasons, the authors questioned if the attempt at formalizing the responsibility of the system designers was at all necessary. The implementation of other measures that have increased the responsibility of designers shows that "formal legislation is only one policy instrument among others and a formal legislation might not even be the most appropriate way to secure a higher degree of responsibility from the system designers" (p. 100). In fact, the government declined a proposal to introduce formal responsibility. The responsibility of system designers still has no other formal basis than the ethical code of conduct developed in Tingvall (1997).

According to McAndrews (2013), however, the effectiveness of relying only on ethical codes is questionable since a code depends on "the experts' self-regulation" and does not generate any leverage for compliance. A study by Van der Burg and Van Gorp (2005) seems to confirm McAndrews's analysis. These authors investigated how engineers involved in the design of trailers understood their moral responsibilities. They found that the engineers' conception of responsibility was limited to the narrow perspective of respecting the traffic laws and designing an economically efficient and physically strong product. They did not seem to consider themselves responsible for finding technological solutions that would improve traffic safety beyond the legal requirements.

As far as we can see, it is not possible to draw any firm conclusions on whether or not a system of accountability for designers of road traffic systems would contribute

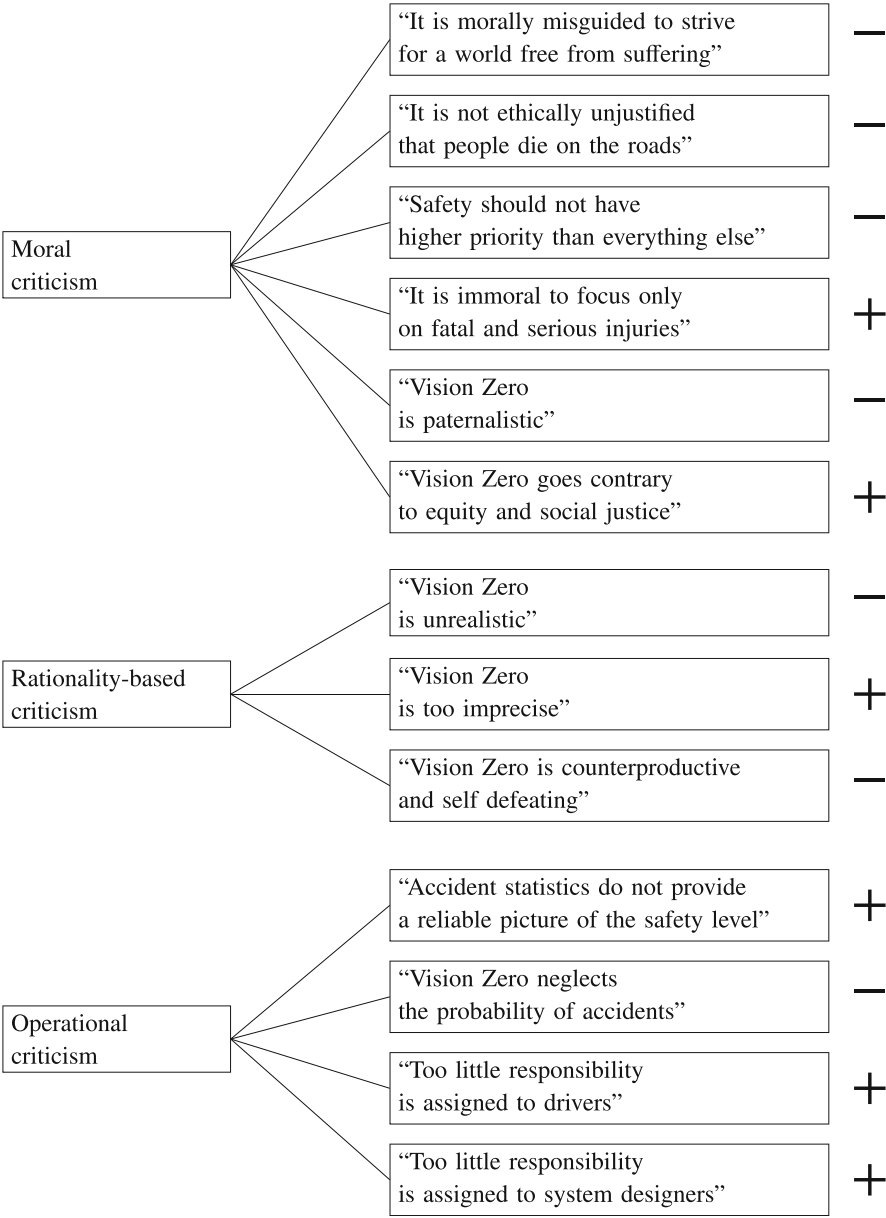


Fig. 3 A summary of our assessments of the arguments discussed in this chapter. The arguments that we found to be useful for a constructive discussion on safety improvements are marked +, whereas the others are marked -

to improved traffic safety. However, the issue is relevant and worth close attention as additional experiences of Vision Zero implementation accumulates. It should definitely be counted as one of the constructive and useful themes of critical discussion.

Conclusion

We have discussed and evaluated 13 arguments. We found that five of them fail because they are based on misrepresentations or misconceptions of Vision Zero (see Fig. 3).

“It is morally misguided to strive for a world free from suffering.” – The goals and ambitions of Vision Zero are much more modest than what these critics claim.

“Safety should not have higher priority than everything else.” – Vision Zero does not include any such claim of absolute priority.

“Vision Zero is paternalistic.” The risk-taking behavior on roads that has to be eliminated according to Vision Zero involves risks for others than the persons who take the risk. Therefore, Vision Zero is not paternalistic.

“Vision Zero is counterproductive and self defeating.” None of the proposed mechanisms that would make Vision Zero counterproductive and self-defeating has been shown to have any impact in practice. Furthermore, the many successes of safety work based on Vision Zero speak against this argument.

“Vision Zero neglects the probability of accidents.” On the contrary, measures that reduce the probability of accidents have a central role in Vision Zero and its implementation.

Two of the arguments are based on correct descriptions of Vision Zero, but they are nevertheless non sequitur arguments:

“It is not ethically unjustified that people die on the roads.” – The proponents of this argument claim that deaths on the roads are acceptable, since people have chosen to risk their lives by travelling on the roads. This argument is fallacious, since most people who are killed on the roads did not wish to take any risks. They just had no other choice than to travel in the risky traffic system that we have.

“Vision Zero is unrealistic.” This criticism is based on a too far-reaching requirement on policy goals. In order for a goal to be rational and useful, it has to be approachable, but it does not necessarily have to be realistic in the sense that it is known beforehand that it can be fully realized. Vision Zero is no doubt approachable to a very high degree.

Finally, we found six of the arguments to be at least in part constructive. They should all be further analyzed and taken into account in future traffic safety work:

“It is immoral to focus only on fatal and serious injuries.” – There are strong moral reasons to give much higher priority to the elimination of fatalities and severe injuries than to the

avoidance of lesser injuries and material damages. However, the critics are right that there is a need to pay more attention to how less serious accidents can be included in safety work that has Vision Zero as its major driving force.

“Vision Zero goes contrary to equity and social justice.” Although this does not apply to Vision Zero in general, the proponents of this argument have been able to show that in some places, Vision Zero activities have increased the burdens of transportation injustices. This is, therefore, a criticism that should be taken seriously and leads to careful evaluations of both procedural and distributive justice in Vision Zero activities.

“Vision Zero is too imprecise.” The critics are right that Vision Zero usually does not come with a precise time plan for what to do and when. It is necessary to complement it with more precise directives and sub-goals, but this has not always been done.

“Accident statistics do not provide a reliable picture of the safety level.” The critics are right that the yearly statistics on deaths in road traffic do not inform us of the risks of rare accidents with many fatalities, such as the collapse of a bridge or a hillside road or a tunnel fire. Traffic safety work based on Vision Zero should pay increased attention to such risks.

“Too little responsibility is assigned to drivers.” Judging by the available evidence, some but not all measures taken to reduce severe accidents can have negative effects on drivers’ sense of responsibility. The risk of such effects should be included in the evaluation of traffic safety measures aiming to implement Vision Zero.

“Too little responsibility is assigned to system designers.” The critics are right that there are currently no means to hold system designers accountable for the design of the road system. It is at present unclear what difference a system of accountability could make or how it should be constructed. However, the issue of accountability should be part of our deliberations on the implementation of Vision Zero.

Cross-References

- [Liberty, Paternalism, and Road Safety](#)
- [Responsibility in Road Traffic](#)
- [Vision Zero and Other Road Safety Targets](#)

References

- Allsop, R. (2005). A long way-but not all the way. Response to the 7th ETSC lecture. Brussels: European Transport Safety Council. http://archive.etsc.eu/documents/7th_Lecture_vision_zero.pdf

- Belin, M.-Å., & Tillgren, P. (2012). Vision zero: How a policy innovation is dashed by interest conflicts, but may prevail in the end. *Scandinavian Journal of Public Administration*, 16(3), 83–102.
- Belin, M.-Å., Tillgren, P., & Vedung, E. (2012). Vision zero – A road safety policy innovation. *International Journal of Injury Control and Safety*, 19(2), 171–179. <https://doi.org/10.1080/17457300.2011.635213>.
- Bergh, T., Carlsson, A., & Larsson, M. (2003). Swedish vision zero experience. *International Journal of Crashworthiness*, 8(2), 159–167. <https://doi.org/10.1533/tjcr.2003.0224>.
- Bullard, R. D. (1990). *Dumping in Dixie: Race, class, and environmental quality* (3rd ed.). Boulder: Westview Press.
- Bullard, R. D., & Wright, B. (Eds.). (2009). *Race, place, and environmental justice after Hurricane Katrina: Struggles to reclaim, rebuild, and revitalize New Orleans and the Gulf coast*. Boulder: Westview Press.
- Centers for Disease Control and Prevention. (1997). Alcohol-related traffic fatalities involving children—United States, 1985–1996. *MMWR: Morbidity and Mortality Weekly Report*, 46(48), 1130–1133.
- Conner, M. (2016). Racial inequity in traffic enforcement. *International Journal of Traffic Safety Innovation*, 1, 14–18. https://www.transalt.org/sites/default/files/201907/VZC_Journal_2016_wCover_0.pdf#page=14.
- Dekker, S. (2017). *The field guide to understanding 'human error'* (3rd ed.). Milton: CRC Press.
- Dekker, S., & Pitzer, C. (2016). Examining the asymptote in safety progress: A literature review. *International Journal of Occupational Safety and Ergonomics*, 22(1), 57–65. <https://doi.org/10.1080/10803548.2015.1112104>.
- Dekker, S. W. A., Long, R., & Wybo, J. L. (2016). Zero vision and a Western salvation narrative. *Safety Science*, 88, 219–223. <https://doi.org/10.1016/j.ssci.2015.11.016>.
- Ecola, L., Popper, S. W., Silbergliitt, R., & Fraade-Blamar, L. (2018). The road to zero: A vision for achieving zero roadway deaths by 2050. *Rand Health Quarterly*, 8(2), 11.
- Edvardsson Björnberg, K. (2008). Utopian goals: Four objections and a cautious defense. *Philosophy in the Contemporary World*, 15(1), 139–154. <https://doi.org/10.5840/pcw200815113>.
- Edvardsson, K., & Hansson, S. O. (2005). When is a goal rational? *Social Choice and Welfare*, 24(2), 343–361. <https://doi.org/10.1007/s00355-003-0309-8>.
- Ekelund, M. (1999). *Varning – livet kan leda till döden!* [Warning – Life can lead to death!]. Stockholm: Timbro.
- Elvebakk, B. (2005). *Ethics and road safety policy*. Oslo: The Institute of Transport Economics (TØI).
- Elvebakk, B. (2007). Vision zero: Remaking road safety. *Mobilities*, 2(3), 425–441. <https://doi.org/10.1080/17450100701597426>.
- Elvebakk, B. (2015). Paternalism and acceptability in road safety work. *Safety Science*, 79, 298–304. <https://doi.org/10.1016/j.ssci.2015.06.013>.
- Elvebakk, B., & Steiro, T. (2009). First principles, second hand: Perceptions and interpretations of Vision Zero in Norway. *Safety Science*, 47(7), 958–966. <https://doi.org/10.1016/j.ssci.2008.10.005>.
- Elvik, R. (1999). Can injury prevention efforts go too far? Reflections on some possible implications of Vision Zero for road accident fatalities. *Accident Analysis and Prevention*, 31(3), 265–286. [https://doi.org/10.1016/S0001-4575\(98\)00079-7](https://doi.org/10.1016/S0001-4575(98)00079-7).
- Elvik, R. (2003). How would setting policy priorities according to cost-benefit analyses affect the provision of road safety? *Accident Analysis and Prevention*, 35(4), 557–570. [https://doi.org/10.1016/S0001-457\(02\)00034-9](https://doi.org/10.1016/S0001-457(02)00034-9).
- Eriksson, L., & Björnskaug, T. (2012). Acceptability of traffic safety measures with personal privacy implications. *Transportation Research Part F: Traffic Psychology and Behaviour*, 15(3), 333–347. <https://doi.org/10.1016/j.trf.2012.02.006>.
- Evans, L. (1996). Comment: The dominant role of driver behavior in traffic safety. *American Journal of Public Health*, 86(6), 784–785. <https://doi.org/10.2105/AJPH.86.6.784>.

- Fugeli, P. (2006). The Zero-vision: Potential side effects of communicating health perfection and zero risk. *Patient Education and Counseling*, 60(3), 267–271. <https://doi.org/10.1016/j.pec.2005.11.002>.
- General Accounting Office. (1983). *Siting of hazardous waste landfills and their correlation with racial and economic status of surrounding communities* (GAO/RCED-83-168). Washington, DC: US. General Accounting Office.
- Global Road Safety Partnership (2007). *Drinking and Driving: a road safety manual for decision-makers and practitioners*. Geneva: International Federation of Red Cross and Red Crescent Societies
- Government Bill 1996/97:137. *Nollvisionen och det trafiksäkra samhället*. Stockholm, Sweden, 22 May 1997.
- Government Offices of Sweden. (2016). *Renewed commitment to Vision Zero. Intensified efforts for transport safety in Sweden*. Stockholm: Government Offices of Sweden.
- Grill, K., & Nihlén Fahlquist, J. (2012). Responsibility, paternalism and alcohol interlocks. *Public Health Ethics*, 5(2), 116–127. <https://doi.org/10.1093/phe/phs015>.
- Hammer, W. (1980). *Product safety management and engineering*. Englewood Cliffs: Prentice-Hall.
- Hansson, S. O. (2007). Philosophical problems in cost–benefit analysis. *Economics and Philosophy*, 23(2), 163–183. <https://doi.org/10.1017/S0266267107001356>.
- Hansson, S. O. (2010). Promoting inherent safety. *Process Safety and Environmental Protection*, 88(3), 168–172. <https://doi.org/10.1016/j.psep.2010.02.003>.
- Hansson, S. O. (2013). The ethics of risk. In *Ethical analysis in an uncertain world* (pp. 118–119). New York: Palgrave Macmillan.
- Hansson, S. O. (2017). Five caveats for risk-risk analysis. *Journal of Risk Research*, 20(8), 984–987. <https://doi.org/10.1080/13669877.2016.1147493>.
- Hansson, S. O., Edvardsson Björnberg, K., & Cantwell, J. (2016). Self-defeating goals. *Dialectica*, 70(4), 491–512. <https://doi.org/10.1111/1746-8361.12161>.
- Harpur, N. F. (1958). Fail-safe structural design. *The Aeronautical Journal*, 62(569), 363–376.
- Hokstad, P., & Vatn, J. (2008). Ethical dilemmas in traffic safety work. *Safety Science*, 46(10), 1435–1449. <https://doi.org/10.1016/j.ssci.2007.10.006>.
- Hoover, K. D. (1990). The logic of causal inference. *Economics and Philosophy*, 6(2), 207–234. <https://doi.org/10.1017/S026626710000122X>.
- Johansson, R. (2009). Vision Zero – Implementing a policy for traffic safety. *Safety Science*, 47(6), 826–831. <https://doi.org/10.1016/j.ssci.2008.10.023>.
- Johnston, I. R., Muir, C., & Howard, E. W. (2014). *Eliminating serious injury and death from road transport: A crisis of complacency*. Boca Raton: CRC Press.
- Jones, M. M., & Bayer, R. (2007). Paternalism & its discontents: Motorcycle helmet laws, libertarian values, and public health. *American Journal of Public Health*, 97(2), 208–217.
- Jones, D. D., Holling, C. S., & Peterman, R. (1975). Fail-Safe vs. Safe-Fail Catastrophes. IIASA Working Paper. WP-75-093. <http://pure.iiasa.ac.at/335/>
- Kerr, S., & LePelley, D. (2013). Stretch goals: Risks, possibilities, and best practices. In E. A. Locke & G. P. Latham (Eds.), *New developments in goal setting and task performance* (pp. 21–31). New York: Routledge.
- Krafft, M., Kullgren, A., Lie, A., & Tingvall, C. (2006). The use of seat belts in cars with smart seat belt reminders – Results of an observational study. *Traffic Injury Prevention*, 7(2), 125–129. <https://doi.org/10.1080/15389580500509278>.
- Kristianssen, A. C., Andersson, R., Belin, M.-Å., & Nilsen, P. (2018). Swedish Vision Zero policies for safety – A comparative policy content analysis. *Safety Science*, 103, 260–269. <https://doi.org/10.1016/j.ssci.2017.11.005>.
- Langeland, T. Å. (2009). *Language and change: An inter-organisational study of the zero vision in the road safety campaign* (Doctoral thesis). University of Stavanger. Retrieved from <http://hdl.handle.net/11250/190896>
- Larsson, P., Dekker, S. W. A., & Tingvall, C. (2010). The need for a systems theory approach to road safety. *Safety Science*, 48(9), 1167–1174. <https://doi.org/10.1016/j.ssci.2009.10.006>.

- Lee, D. J. (2018). Delivering justice: Food delivery cyclists in New York City. CUNY Academic Works. https://academicworks.cuny.edu/gc_etds/2794
- Lind, G., & Schmidt, K. (2000). *Leder nollvisionen till det trafiksäkra samhället?: en kritisk analys av effektiviteten i trafiksäkerhetsarbetet*. Stockholm: Kommunikationsforskningsberedningen (KFB).
- Locke, E. A., & Latham, G. P. (2002). Building a practically useful theory of goal setting and task performance. *American Psychologist*, 57, 705–717. <https://doi.org/10.1037/0003-066X.57.9.705>.
- Long, R. (2012). *The zero aspiration, the maintenance of a dangerous idea*. Kambach: Human Dimensions.
- Lugo, A. (2015, September 30). Unsolicited Advice for Vision Zero. Urban Adonia. Retrieved from <http://www.urbanadonia.com/2015/09/unsolicited-advice-for-vision-zero.html>
- Lutter, R., & Morrall, J. F., III. (1994). Health-health analysis: A new way to evaluate health and safety regulation. *Journal of Risk and Uncertainty*, 8, 43–66. <https://doi-org.focus.lib.kth.se/10.1007/BF01064085>.
- McAndrews, C. (2013). Road safety as a shared responsibility and a public problem in Swedish road safety policy. *Science, Technology and Human Values*, 38(6), 749–772. <https://doi.org/10.1177/0162243913493675>.
- McKenna, F. P. (2007). The perceived legitimacy of intervention: A key feature for road safety. In *Improving traffic safety culture in the United States: The journey forward* (pp. 165–175). Washington, DC: AAA Foundation for Traffic Safety.
- Mendoza, A. E., Wybourn, C. A., Mendoza, M. A., Cruz, M. J., Juillard, C. J., & Dicker, R. A. (2017). The worldwide approach to vision zero: Implementing road safety strategies to eliminate traffic-related fatalities. *Current Trauma Reports*, 3, 104–110. <https://doi.org/10.1007/s40719-017-0085-z>.
- Morgan, R. (2018). *Safe system or Stalinist system? Road safety at any cost*. Bulleen: Robert Morgan Traffic Engineering.
- Morse, S. (2015). Policing and Safe Streets Where are the Planners? Progressive Planning Magazine. NO. 202 | WINTER 2015. ISSN 1559-9736 http://www.plannersnetwork.org/wp-content/uploads/2015/02/PPM_Winter2015_WEB.pdf
- Moskowitz, G., & Grant, H., (eds.) (2009). *The Psychology of Goals*. The Guilford Press, New York.
- Neumayer, E., & Plümper, T. (2016). Inequalities of income and inequalities of longevity: A cross-country study. *American Journal of Public Health*, 106(1), 160–165. <https://doi.org/10.2105/AJPH.2015.302849>.
- Nihlén Fahlquist, J. (2009). Saving lives in road traffic – Ethical aspects. *Journal of Public Health*, 17, 385–394. <https://doi.org/10.1007/s10389-009-0264-7>.
- Perrow, C. (1999). *Normal accidents*. Princeton: Princeton University Press.
- Rall, M., Manser, T., Guggenberger, H., & Gaba, D. M. (2001). Patient safety and errors in medicine: Development, prevention and analyses of incidents. *Anesthesiologie Intensivmedizin Notfallmedizin Schmerztherapie*, 36(6), 321–330. <https://doi.org/10.1055/s-2001-14806>.
- Reason, J. (2000). Safety paradoxes and safety culture. *Injury Control and Safety Promotion*, 7(1), 3–14. [https://doi.org/10.1076/1566-0974\(200003\)7:1;1-V;FT003](https://doi.org/10.1076/1566-0974(200003)7:1;1-V;FT003).
- Roos, A., & Nyberg, J. (2005). *Kommunpolitikens syn på trafiksäkerhet: En intervjustudie* (VTI rapport 527, with a summary in English). Linköping: VTI.
- Rosencrantz, H., Edvardsson, K., & Hansson, S. O. (2007). Vision zero – Is it irrational? *Transportation Research Part A: Policy and Practice*, 41(6), 559–567. <https://doi.org/10.1016/j.tra.2006.11.002>.
- Schlosberg, D. (2007). *Defining environmental justice: Theories, movements, and nature*. Oxford: Oxford University Press.
- Sherratt, F., & Dainty, A. R. J. (2017). UK construction safety: A zero paradox? *Policy and Practice in Health and Safety*, 15(2), 108–116. <https://doi.org/10.1080/14773996.2017.1305040>.
- Sitkin, S. B., See, K. E., Miller, C. C., Lawless, M. W., & Carton, A. M. (2011). The paradox of stretch goals: Organizations in pursuit of the seemingly impossible. *Academy of Management Review*, 36(3), 544–566. <https://doi.org/10.5465/amr.2008.0038>.
- Sveriges Bussföretag. (2018). Statistik om bussbranschen Mars 2018. Downloaded March 29, 2020 from <https://www.transportforetagen.se/globalassets/rapporter/buss/statistik-om-bussbranschen-2018-03.pdf?ts=8d79ff684808180>

- Tingvall, C. (1997). The zero vision: A road transport system free from serious health losses. In H. von Holst & Å. Nygren (Eds.), *Transportation, traffic safety and health* (pp. 37–57). Berlin/Heidelberg: Springer. https://doi.org/10.1007/978-3-662-03409-5_4.
- Tingvall, C., & Haworth, N. (1999). Vision Zero—An ethical approach to safety and mobility, paper presented to the 6th international conference road safety & traffic enforcement: beyond 2000.
- Transport Styrelsen. (2020). Trafiksäkerheten i Sverige 2019 Preliminär statistik över omkomna och allvarligt skadade i vägtrafik, spårtrafik, sjöfart och luftfart. Norrköping. https://www.transportstyrelsen.se/globalassets/global/om_oss/trafiksakerheten-i-sverige/trafiksakerheten-i-sverige-preliminar-olycksstatistik-2019.pdf
- United Church of Christ. Commission for Racial Justice. (1987). Toxic wastes and race in the United States: A national report on the racial and socio-economic characteristics of communities with hazardous waste sites. Public Data Access.
- Van der Burg, S., & Van Gorp, A. (2005). Understanding moral responsibility in the design of trailers. *Science and Engineering Ethics*, 11(2), 235–256.
- Viscusi, W. K., Magat, W. A., & Huber, J. (1991). Pricing environmental health risks: Survey assessments of risk-risk and risk-dollar trade-offs for chronic bronchitis. *Journal of Environmental Economics and Management*, 21, 32–51. [https://doi.org/10.1016/0095-0696\(91\)90003-2](https://doi.org/10.1016/0095-0696(91)90003-2).
- Vision Zero Equity Strategies for Practitioners, Vision Zero Network, February 13, 2018. http://visionzeronetwork.org/wp-content/uploads/2017/05/VisionZero_Equity.pdf
- Vissers, L. H. (2017). *Alcohol-related road casualties in official crash statistics*. Paris: International Transport Forum ITF, 2017, 55 p., ref.; International Traffic Safety Data and Analysis Group IRTAD Research Report.
- World Health Organization. (2018). *Global status report on road safety 2018*. Geneva: World Health Organization.
- Young, I. (1990). *Justice and the politics of difference*. Princeton: Princeton University Press.
- Zwetsloot, G. I. J. M., Aaltonen, M., Wybo, J.-L., Saari, J., Kines, P., & Op De Beeck, R. (2013). The case for research into the zero accident vision. *Safety Science*, 58, 41–48. <https://doi.org/10.1016/j.ssci.2013.01.026>.
- Zwetsloot, G. I. J. M., Kines, P., Wybo, J.-L., Ruotsala, R., Drupsteen, L., & Bezemer, R. A. (2017). Zero Accident Vision based strategies in organisations: Innovative perspectives. *Safety Science*, 91, 260–268. <https://doi.org/10.1016/j.ssci.2016.08.01>.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.



What Is a Vision Zero Policy? Lessons from a Multi-sectoral Perspective

4

Ann-Catrin Kristianssen and Ragnar Andersson

Contents

Introduction	153
Theoretical and Analytical Framework	154
Governing by Visions	154
Turning Visions into Policies and Goals	155
Policy Content as a Framework for Description and Analysis	157
The Five Vision Zero Policies	158
Vision Zero for Road Traffic Safety	158
Vision Zero for Fire Safety	160
Vision Zero for Patient Safety	161
The Vision Zero for Suicide	162
The Vision Zero for Workplace Safety	162
Analyzing Differences and Similarities in Vision Zero Components	164
Problem Framing	164
Monitoring and Surveillance	165
Means, Programs, and Governing Structures	167
Discussion	169
Are There Discernable Determinants for a Vision Zero and for It Being Successful?	169
Building the Vision Zero Ship at Sea?	171
A Conceptual Distinction	171
Cross-References	172
References	172

A.-C. Kristianssen (✉)

School of Humanities, Education and Social Sciences, Örebro University, Örebro, Sweden
e-mail: ann-catrin.kristianssen@oru.se

R. Andersson

The Centre for Societal Risk Research, CSR, Karlstad University, Karlstad, Sweden
e-mail: ragnar.andersson@kau.se

© The Author(s) 2023

K. Edvardsson Björnberg et al. (eds.), *The Vision Zero Handbook*,
https://doi.org/10.1007/978-3-030-76505-7_4

151

Abstract

Vision Zero is a term mainly connected with road traffic safety and has its roots in the Swedish road safety strategy. It was formally adopted by the Swedish parliament in 1997, and due to the initial success of lowering the number of deaths in traffic crashes significantly, it has become a role model for road safety strategies in countries and cities all over the world. In Sweden, Vision Zero for road safety has also inspired the introduction of Vision Zero policies in other sectors, and this chapter focuses on Vision Zero from a multi-sectoral perspective. The purpose of this chapter is twofold: to present five different cases of Vision Zero policies and to discuss what constitutes a Vision Zero policy based on these five cases. The five cases are found in road traffic safety, fire safety, patient safety, suicide, and workplace safety. Every case has its unique preconditions in terms of laws, actors, scope, etc., but they are also similar in relation to injury prevention and the ambition to decrease the number of deaths and serious injuries. The five Vision Zero policies are summarized by presenting the problem and problem framing, the goal, measures, and solutions as well as leading actors and governing structures. We find that the problem itself is quite self-explanatory in each case but that the problem framing and attribution of responsibility differ. All cases have on paper been inspired by the road safety strategies, but the systems approach, so intimately connected with Vision Zero, is more or less absent in the cases of fire safety and suicide. Furthermore, in the field of fire safety, responsibility is placed on the individual and on the business sector rather than based on a shared responsibility and ultimately on the system designers. In all five cases, there are a set of measures in place, but there are differences in implementation due to temporal factors and also what kind of governing and steering structures are in place. There is also a difference in internal support where the Vision Zero for suicide stands out as having less support among agencies working with the issue. Finally, the monitoring systems differ from case to case. The Vision Zero for road traffic safety stands out as having a monitoring and evaluating system based on specific safety targets ultimately aiming toward zero (management by objectives). Based on the empirical findings, we argue that besides having a clear problem and problem framing, a toolbox of measures, a monitoring system, and a governing structure, a policy based on a visionary approach with an ambition to reach zero needs additional perspectives or criteria in order to be successful: (1) a **scientific** approach to problem framing and solutions, (2) a **comprehensive** approach, (3) a **long-term** commitment, and (4) a system and structure based on **governance**. These criteria do not necessarily have to be in place in order to adopt Vision Zero, but they are a prerequisite for building a system based on Vision Zero.

Keywords

Vision Zero policies · Comparative policy content analysis · Road traffic safety · Suicide · Workplace safety · Fire safety · Patient safety

Introduction

The term Vision Zero has received broad international attention after its launch in Sweden in the end of the 1990s as a road safety measure (c.f. Belin et al. 2012; Belin and Tillgren 2012). Vision Zero for road traffic safety contains a number of ethical and practice-oriented approaches and measures. The term and at least parts of the policy package of Vision Zero have diffused to other countries, cities, and policy areas (c.f. Elvebakk 2007). A quick search on the Internet gives plenty of results from all over the world. Vision Zero is not only discussed in terms of road safety, but as a vision that can be implemented in many policy areas. In Sweden, there are Vision Zero policies in many sectors, such as within healthcare, and there are ongoing discussions about introducing Vision Zero policies in yet more areas, for instance, drowning. Vision Zero approaches are adopted in various spheres, such as public administration, private companies, and organizations.

The questions raised in this chapter depart from an empirical material where five Vision Zero policies within various policy areas in Sweden related to injury prevention are compared. Parts of the empirical material were analyzed from a different perspective in an article from 2018 (Kristianssen et al. 2018). The areas in question are road safety, workplace safety, patient safety, fire safety, and suicide. The Vision Zero for road safety is, in the documents analyzed in the article, described as a role model for the other Vision Zero policies. The policies have at least four elements in common: they are all nationally adopted, they are adopted within one single country (Sweden), they are all related to injuries (a medically reasonably homogenous and well-defined area), and they all present a vision for zero fatalities. What differs is that they are applied in different policy areas which entail variances in the specific preconditions of each policy area, both in terms of actors and structures.

Departing from that empirical material, the purpose of this chapter is twofold: first, to provide a short description of each Vision Zero policy and, second, to scrutinize and discuss what characterizes Vision Zero from a conceptual point of view and what it should contain from a normative perspective to be able to manage a zero approach. This is important because visions are often used directly as, or transformed into, steering tools. The question is if it is always appropriate to use a vision as a tool for implementing policy. It entails both opportunities and risks, and the zero approach has been labeled both paternalistic (Ekelund 1999; Elvebakk 2015), inefficient (Elvik 2003), and unrealistic (Lind and Schmidt 1999). On the other hand, research has also shown that the innovative approach of Vision Zero enables actors and institutions to break away from habits and known patterns of behavior (Belin et al. 2012) and that zero fatalities and serious injuries are not an irrational goal from a conceptual point of view (Rosencrantz et al. 2006). The conceptual perspective of Vision Zero is furthermore important to study as there are several similar approaches, but with clear differences, such as the concept of zero tolerance.

The chapter is divided into five parts starting with this introduction. The second part will explore the use of visions as a steering and governing tool and how visions

relate to policy. The third section contains a description of the five Vision Zero policies, and the fourth part is a comparative analysis. The fifth and final section of the chapter is devoted to a discussion about the concept of Vision Zero, related terms, as well as to principles regarding governing by Vision Zero.

Theoretical and Analytical Framework

Controlling adverse occurrences, whether environmental, social, or health-related as in the case of injuries, often constitutes considerable challenges to modern societies since their determinants are rarely confined to single policy domains in reach of traditional top-down governmental initiatives. On the contrary, such problems, often referred to as wicked problems (Rittel and Webber 1973), are complex by nature with determinants rooted across a spectrum of policy areas and sectors, calling for broader collaborative approaches, often referred to as “governance” (c.f. Hedlund and Montin 2009). These problems require a broad set of measures, subject to continuous scrutiny and evaluation, as well as a long-term commitment to finding solutions to the problem. There are a number of models and practices available when working with solutions for these kinds of broader societal problems. They stem from various policy fields and perspectives such as foresight or backcasting models, policy innovation, reforms, strategy management, mission statements, and steering by vision. Research on these topics is performed within many disciplines such as future studies, policy studies, studies in technology development, engineering, business and management studies, etc.

New solutions to major societal problems are thus wrapped in different terms, but the intention to change the current situation and create a different future is the same. Some approaches are focusing on reforming a current system, not necessarily creating a completely new one. Other approaches embrace the idea of a more or less complete overhaul or replacement of the old system. In public administration, reforms are common whether they target small changes or larger transformations. But reforms rarely replace earlier changes, which lead to a layered system (Christensen and Lægreid 2012). This kind of fragmentation has prompted reforms focusing on governance and coordination (c.f. Pollitt 2003). Reforms and policy changes are thus a common part of the everyday routines in both public administration and the private sector, often targeting a specific problem or issue.

Governing by Visions

As mentioned, there are specific tools today for working with comprehensive and complex societal problems, and the ambition of these tools is to have a much broader transformative aim. We will discuss a few of these terms and approaches with a focus on visions as a policy tool. One of those is the use of strategies. A strategy is an “...engine of change, a mechanism to transform the present and mold it in the image of a desired future to come” (Kornberger 2013). Strategies are used in politics, in

public administration, in the business sector, and so on, and some are more short term and changeable, while others are long term by nature. The long-term capability of strategies makes it possible to transcend spatial boundaries, temporal restraints, and current challenges. A strategy creates a vision of what the ideal future could look like.

Related to strategies, it is becoming increasingly more common to use visions as a policy and governing tool. In earlier research, visions were often referred to as mission statements (Weiss and Piderit 1999). Using visions as a tool to change the present system entails both opportunities and problems. A vision is intended to inspire and to make people consider new approaches and methods (Hallström and Grafström 2016). A vision does not normally include detailed measures, and that provides flexible opportunities for actors to adapt to the vision (Gioia et al. 2012). A vision is also a long-term commitment with a core message or core image of what the future should look like. On the other hand, as a vision tends to be very broad, the risk is that it turns into nothing but beautiful words. If there is no substance to a vision in the sense that it is translated into specific methods and measures used in order to reach the goal of the vision, then the vision can still inspire, but will not necessarily lead to the intended result. Furthermore, visions sometimes tend to be exclusive rather than inclusive because of the focus on reaching the end goal. Alternative visions or paths to reach a certain goal are excluded in the narrative of the vision (Dignum et al. 2018). On the other hand, for a vision to work, it has to be interpreted and implemented in a comprehensive way reaching as many sectors, actors, and aspects as possible. So, there is duality as the vision needs to be both focused and exclusive in order to focus on and reach a specific goal and at the same time inclusive enough to convince and involve as many relevant aspects and actors as possible. In other words, there has to be dedicated actors driving the transformation process, and the actual problem to be solved has to be well-defined and accepted as a societal problem that needs to be handled with a long-term visionary approach. Dignum et al. (2018) describe the performativity of a vision related to its scope, i.e., reaching beyond what was earlier possible, to its systemic and holistic capability, to its problem description and to its description of the values threatened in the current system, and to its inclusion of a framework for targets and for monitoring the progress. Visions can be studied from various perspectives: (1) as a process, (2) as content, and (3) as output (Dignum et al. 2018). This chapter will concentrate on the content perspective.

Turning Visions into Policies and Goals

The use of visions provides an opportunity to inspire and to think creatively about a societal problem, but there is still a need to present a realistic plan on how to reach that vision. Therefore, we often see a vision complemented by a more specific plan or program for implementation. A vision turned into a more concrete policy program can sometimes be seen or presented as a policy innovation (Belin et al. 2012; Sørensen 2016) where the policy does not only contain specific measures and

solutions but clearly draws from visionary aspects for the future. “The content of a policy is innovative to the extent that it offers a new definition of a political problem, provides a new political vision for the political community, and/or proposes a new set of political goals and strategies” (Sørensen 2016:157). The uniqueness of a policy innovation is that it departs from visions and strategies but uses a specific method or approach to work with those goals and strategies.

One such line of methods are goal-related approaches, and one example is management by objectives (MBO), sometimes also called management by results (MBR). This approach was first introduced in the private sector to promote productivity, but came to influence the public sector as well, and was soon incorporated as a core component in what later became known as the New Public Management (NPM) approach (Hood 1991; Lægreid 2011). This development was also driven by a political ambition to save public expenditures by means of privatization and competition among providers. The role of politicians should be restricted to clarifying needs, setting goals, and reviewing results, according to the proponents, while meeting the goals should preferably be left to private providers in competition, when possible and appropriate. In addition, the particular field of accident and injury prevention, especially in industrial settings, was also influenced by parallel industrial developments in quality control, such as quality assurance, total quality management, and the like (Kjellén 2000). The core idea was that undesired outcomes were better controlled by means of proactive identification and control of upstream deviations and determinants, instead of dealing with problems in retrospect. This in turn presupposes a thorough understanding of the underlying causes of adverse outcomes and access to valid measures thereof. These developments found their way to broad public applications as well. Working toward specific goals in relation to a complex societal problem has led to an understanding that there is often a need for broader partnerships between key actors in order to reach that goal.

The zero approach is a vision and a goal with an innovative intent and ethical core. Getting the numbers down to zero, whether it concerns domestic abuse or traffic crashes, is for many a reasonable goal. For some actors it is also a question related to morality and ethics, asking whether it is morally or ethically acceptable that people die due to injuries in, for instance, traffic crashes or fires. The zero approach can also be connected to the so-called improvement principles that are based on some kind of zero perspective although these principles cannot necessarily be incorporated into a model such as Vision Zero (Hansson 2019). As Vision Zero diffuses all over the world, to different levels and to various sectors, it is important to distinguish Vision Zero from other zero perspectives. Based on the summary of each Vision Zero policy, we will in the last section of this chapter discuss and make a distinction between Vision Zero and other zero perspectives. We will use zero tolerance as an example of a zero perspective, but with basically the opposite approach to human behavior. A visionary approach and the concept of zero can be interpreted as a reasonable combination, as both concepts concern long-term commitments. Just as with other visions, the challenge is how to transform Vision Zero into a workable policy tool. The implementation phase is thus crucial for the success or failure of the approach. The Vision Zero for road traffic safety is a policy program

targeting a growing societal problem with specific long-term scientific, ethical, and administrative approaches, which we will return to in the next section.

Policy Content as a Framework for Description and Analysis

The next section contains a summary of the earlier mentioned Vision Zero policies. The content of each Vision Zero policy will be summarized using four categories inspired by a model presented in Kristianssen et al. (2018). Although using in part a similar model, this chapter contains a new analysis based on a different purpose and theoretical approach. The analytical framework is also complemented by the earlier mentioned performative aspects of a visionary approach (Dignum et al. 2018). First, what is the **problem** to be solved in each policy area; second, what is the **goal**; third, what **measures and solutions** will solve the problem; and fourth, what **actors** will solve it. More specifically, in order to solve a major societal problem, it is necessary to understand the nature of the problem, its determinants, and its scope, diffusion, and development within all societal sectors. It is furthermore necessary to have access to a credible toolbox containing measures and solutions for how to deal with the problem. This entails a robust program for systematic implementation and evaluation. These aspects related to problem and measures are dependent on the actual goal and governing structures. As mentioned, goals can be set up in different ways with different ambitions, and the measures and solutions are dependent on that ambition. What actors are involved in deciding, prioritizing, and implementing measures can lead to different results. Using a governance structure with the active involvement of broad networks covering many relevant actors within a specific field can from one perspective alleviate the implementation of important measures and from another create a more fragmented implementation process. Having a strict central steering process risks excluding important actors, but the advantage could be more efficient and/or faster processes.

The descriptions will help us to understand the basic components of each Vision Zero policy. It will furthermore provide information on the conceptual development of Vision Zero, and departing from the descriptive findings, we will analyze the five cases comparatively by using the following questions;

1. What is the scope of the problem framing in each of the five cases and are the problems presented scientifically determined?
2. Is there a functioning monitoring system related to each case and does it have a long-term transformation focus?
3. Do the measures and solutions form a comprehensive policy program for each case with a designated governing structure related to each Vision Zero and if so, is the governing structure centralized or network based?

The interpretation and development of what constitutes a Vision Zero is particularly interesting in the light of its diffusion from road traffic to other sectors. The concept is challenged by the differences in preconditions, and any claims of a

generalization of the concept are on the line. In the Discussion section, we will use the findings from the summary and comparative analysis to:

1. Approach the concept of Vision Zero by discussing whether there are discernable determinants for a Vision Zero regardless of policy area and for a successful implementation
2. Discuss whether these determinants have to be in place before adopting a Vision Zero or if structures and system can be created afterward. In other words, is it possible to build a ship while at sea?
3. Distinguish the boundaries between Vision Zero and other zero perspectives

The Five Vision Zero Policies

The descriptions below are based on policy documents related to the actual adoption of the Vision Zero policies and the period leading up to the adoption. We will describe Vision Zero for road safety in more detail as it served as the role model for the other Vision Zero policies described in this chapter. The time period for road safety strategies stretches a bit longer as it served as an example for at least 5–10 years.

Vision Zero for Road Traffic Safety

The Vision Zero for road traffic safety was adopted in 1997 by a parliament decision (Swedish Parliament 1997b; see also Swedish Parliament 1997a; Swedish Government 1997). The decision stated that “no one shall die or be seriously injured in road traffic.” The vision was furthermore underpinned by supporting theories regarding the problem description, ethical and strategic perspectives, and a steering and implementation model related to scientific evidence. The intention of the theoretical support was to show a credible policy package aimed at systematically reducing the number of deaths and serious injuries over time.

The adoption of the vision has been described as a paradigm shift in road traffic safety (Tingvall and Haworth 1999; Belin et al. 2012). The most important changes relate to the responsibility of the individual in relation to the responsibility of the system designer (Nihlén Fahlqvist 2006) and what the problem at hand is – the crash or the injury. The previous road safety work was based on the so-called human factor approach, according to which it is the responsibility of the individual to avoid accidents. Vision Zero on the contrary focuses on shared responsibility, where there is a complementary responsibility with the system designers (i.e., road design, vehicle design, etc.), and highlights prevention of injuries rather than prevention of accidents. Injuries are regarded as the major problem, particularly deaths and serious injuries, while damages to properties must be tolerated as they are necessary to shield the human being from injuries. Another key aspect is that human mistakes have to be

taken into account in designing the system, as it is part of the human nature to make mistakes.

The ethical foundation is that deaths and serious injuries are not accepted within the transport system. The alternative to tolerate a certain number of deaths or serious injuries is not seen as an option in comparison. There may be a way to calculate a balance between the cost of injuries and the benefit of mobility using cost-benefit analyses, but the development of new technology has a tendency to challenge this balance, by creating new ways of prevention. Therefore, the only reasonable conclusion is to strive for zero even though it might take time. There is a strong connection between this long-term approach and terms that were launched regarding quality development already during the 1970s and 1980s such as “continuous improvements” and the like. These terms have been used in both the private and public sector in order to systematically increase quality.

The scientific or evidence-based approach of Vision Zero is to base the road safety work on scientific results as well as on successful policies and approaches. Injury prevention has since the end of World War II been connected to preventive medicine and public health related to a general prevention of health problems. It concerns preventing or limiting harmful exposure that can be sudden or long term, or counteracting the consequences of such an exposure through protection, rescue, care, or rehabilitation. The details differ depending on what specific risk we are looking at, but the principles are more or less commonly applicable. It means that there is a solid scientific base to rely on when it comes to understanding the preconditions for preventing deaths and serious injury even within a policy area such as road safety. Applied to injuries related to sudden events, there are a number of chronological aspects: to prevent or minimize the negative effect of the event itself, to stop or limit the negative consequences of the injury (the consequence of the event), and to take care of, treat, and rehabilitate the injury. Significant efforts have been made to lessen the consequences of road traffic crashes, often with great success. Vehicles are safer, barriers prevent vehicles from crashing off the road, and poles are folding when hit. These are all safety interventions made to shield the human being from being exposed to potential lethal violence. Speed limits reduce the potential violence but also lessen the risk for crashes by increasing the driver’s control of the vehicle. The number and scope of possible measures are comprehensive, and the technological development is constantly producing new possibilities. It is ultimately the responsibility of the system owner to craft the modern and safe transport system design in relation to mobility demands, environmental concerns, and accessibility. The main principle of Vision Zero is that human tolerance for crash violence (Haddon 1968) should guide the design of the transport system.

Finally, the steering model for road traffic rests upon a safe system design where several actors are viewed as system designers. There are numerous actors involved in particularly implementing road safety measures, which in some ways has led to a fragmented governing system. The Swedish Transport Administration has a lead role today in the development of road safety measures, but this has not always been the case during the last 10 years. In order to find an efficient way to work with road safety, several of structures have been set up. Networking is one method to bring

actors together. There are various networks discussing and analyzing the current status of road safety in Sweden, a number of them led by the Swedish Transport Administration. Studies show that these networks tend to be set up more for information exchange than focusing on decision-making or raising public awareness (Hysing 2019). Another method used to coordinate efforts is the system of management by objectives. This means continuous measuring of the number of deaths and serious injuries and identifying and monitoring the most important indicators of road safety over time and also giving feedback to the relevant actors.

Vision Zero for Fire Safety

The Vision Zero for fire safety was launched in 2010 by the Swedish Civil Contingencies Agency (MSB). The vision, stated in the national guidelines for fire safety, said that: “no one shall die or be seriously injured due to fire in Sweden” (Swedish Civil Contingencies Agency 2010: 5, our translation). Prior to the adoption of the Vision Zero, national strategies had been discussed for a long time leading to updates in laws and regulations (c.f. Swedish Government 2002). The main problem is that around 100 individuals die every year due to fire and approximately 1000 individuals are seriously injured. The new laws that had been adopted did not lead to a reduction in these numbers, which prompted the introduction of a Vision Zero. According to the national guidelines, the responsibility for fires is placed on the individual and on the business sector, although there is an awareness of the theories regarding human errors and the limitations of placing responsibility solely on the individual. This is a clear break from the systems approach presented in the Vision Zero for road traffic. One explanation is that existing policy and legal frameworks at times create obstacles for implementing measures related to a systemic perspective.

The ultimate goal of zero fatalities is intended to be reached by using interim goals. The results from each time period will be thoroughly evaluated. The first period stretches to 2020. The national guidelines present a number of measures divided into four strategic areas, knowledge and communication, technological solutions, local coordination and collaboration, and evaluation and research. In each of these areas, separate measures were presented such as scrutinizing databases, campaigns, increased collaboration, specific technical innovations, etc. (Swedish Civil Contingencies Agency 2010). These measures were not linked together theoretically in the guidelines, and no coherent steering or governing model was presented to clarify or develop the issue of responsibility to make credible the long-term abilities of the vision.

The Swedish Civil Contingencies Agency has a lead role in evaluating the national guidelines, but the implementation of the Vision Zero falls on several actors, particularly the local authorities. There is a national advisory committee including members from all kinds of societal institutions. But, issues of responsibility and steering are leaning very much on the law regulating fire safety in Sweden (Swedish Law on accident prevention 2003) which places the main responsibility for fire safety on the individual. As the law tends to limit the scope of Vision Zero, there are

a number of initiatives today focusing on outlining a system's approach for the area of fire safety (see the ► [Chap. 38, "Vision Zero on Fire Safety"](#)).

Vision Zero for Patient Safety

The Vision Zero for patient safety was presented in 2013 in a national strategy produced by the National Board of Health and Welfare. This was an assignment from the government and the document stated that "Vision Zero is the image of a future where human beings do not die or are seriously injured within the health or dental care system" (The National Board of Health and Welfare 2013: 8, our translation). The Vision Zero for patient safety was preceded by a number of discussions and initiatives, such as the introduction of a new law on patient safety from 2010 (Swedish Law on Patient Safety 2010: 659; Swedish Government 2007; Swedish Government Official Investigations 2008; Swedish Government 2009).

The main problem presented in the national strategy is that 100,000 individuals are injured every year in the Swedish healthcare system, which means approximately 9% of all patients treated at hospitals. The main reasons for these injuries are poor routines, that regulations are not followed, failures in leadership, and regional differences. The responsibility for reducing injuries and deaths fall on the healthcare system (the National Board of Health and Welfare 2013). The national strategy presents a clear system's approach and the writings have been inspired by the ideas and theories regarding the human factor presented in the Vision Zero for road safety.

In order to reach zero, effect goals have been introduced. They focus on patient safety culture, increasing patient participation, reducing the number of frequent and serious healthcare injuries, and increasing knowledge about effective measures and when to implement these measures. The vision consists of a list of 16 areas where measures are needed but no coherent theoretical framework is presented. On the other hand, the steering model puts the caregiver as responsible for the system, and there is a plan for systematic safety work in line with a continuous improvement approach.

The National Board of Health and Welfare is the lead agency in the sense that it produces reports and acts as a coordinator, but many actors are working within this area. You can find them in both the private and the public sector. One particularly important actor is the Health and Social Care Inspectorate that deals with complaints and irregularities in the healthcare system based on the Patient Safety Law. The inspectorate produces reports and statements regarding the state of the Swedish healthcare system. The vast number of actors working with patients makes it difficult to grasp what is the system to be monitored in accordance with Vision Zero. In addition, Sweden has a considerable number of private actors within the health sector as well as actors on different levels of public administration, which leads to a challenge concerning coordination. One risk and sometimes also a direct consequence are differences in quality, methods, and techniques depending on region, which can make healthcare geographically unequal.

The Vision Zero for Suicide

The Vision Zero for suicide was decided by the Swedish parliament in 2008 after a government proposal. The vision states that “no one should find him- or herself in such an exposed situation that the only conceivable way out is suicide. The government’s vision is that no one should have to end their life” (Swedish Government 2007, our translation). The decision to adopt a Vision Zero for suicide was not based on clear requests from actors working with suicide. On the contrary, the national strategy produced by the National Board of Health and Welfare and the Public Health Agency in 2006 argued against a Vision Zero policy. “The ethical problems related to suicide prevention cannot be completely solved. Therefore, it is not appropriate to formulate a Vision Zero for suicide similar to the Vision Zero for road traffic fatalities. It is possible though to work towards reducing the number of suicides.” (National Board of Health and Welfare and the Public Health Agency 2006: 27, our translation).

The main problem to be solved is that approximately 1500 individuals commit suicide every year. For a long time, this was seen as an individual problem, and one reason for this is that every suicide is complex and cannot be generalized. Vision Zero for suicide added parts of a system’s approach to the policy area by placing the responsibility for suicide prevention on the healthcare system and its work to identify and support individuals in risk of committing suicide (Swedish Government 2007). Also in relation to this policy area, the government has been greatly inspired by the Vision Zero for road traffic safety.

In conjunction with the Vision Zero decision in 2008, the parliament adopted a nine-point program for suicide prevention, which was a mix of both measures and effect goals. These focus on the production of information material particularly for school pupils, on reducing alcohol consumption, on reducing access to lethal means in all kinds of societal contexts, on creating a national function for knowledge assessment, on continuing preventive work within the healthcare system, on gathering and analyzing research results, on initiating campaigns, on improving statistics, and on supporting voluntary organizations in their suicide preventive work (Swedish Government 2007).

These measures and goals did not present a coherent theoretical framework. The parliament decision does not reveal a clear steering and governing model or implementation scheme regarding how the vision will be carried out or who is responsible for what. The Public Health Agency and the National Board of Health and Welfare are key actors in implementing the vision, for setting up measures, and for gathering knowledge and spreading information, but issues of steering are still an ongoing discussion within this field. The critique is still widespread (Tryssel 2018), and the number of actors and levels that constitute the system is considerable, just as in the case of patient safety.

The Vision Zero for Workplace Safety

The Vision Zero for workplace safety was presented in a government proposal in 2016. It says that “No one should have to die as a result of their job. Concrete measures are necessary in order to prevent work-related accidents leading to

injury or death” (Swedish Government 2016a, our translation). This was preceded by growing concerns that not all accidents were reported and that workplace safety was becoming more fragmented and harder to monitor. The Swedish parliament therefore urged the government in 2014 to initiate a dialogue concerning fatal accidents, to encourage more research and education within this field, to improve statistics, and to reduce the number of workplace-related incidents such as bullying. The Vision Zero for fatal accidents was one of three parts of the government strategy from 2016 and the other areas focused on a sustainable working life and psychosocial work environment (Swedish Government 2016a, b, c).

The main problems presented by various actors in the field as well as in the government decision were the growing number of accidents in the workplace, but not necessarily fatal accidents. There were also growing concerns regarding the upward trend of longer periods of sick leave. Another problem was the large number of actors with workplace activities in Sweden, making it hard to monitor workplace safety. One cause of these problems was related to the fact that there are more foreign entrepreneurs active in Sweden, not necessarily following or having the knowledge of Swedish workplace law. Another cause is the growing number of short-term employments, increasing migration and movement of people, and more sub-entrepreneurs. These are, for different reasons, risk factors in terms of safety (Swedish Government 2016a).

In relation to the government strategy from 2016, a number of investigations were launched but no long-term strategy for the realization of the vision was presented. The Swedish Work Environment Authority is the lead agency concerning analyzing and evaluating the development of the policy area. The Authority has been given the task to increase supervision and monitoring of both Swedish and foreign companies. The work will also include a gender perspective as well as improvements regarding information and communication. The steering model for workplace safety is called systematic safety work (SAM) (the Swedish Work Environment Authority 2001) and has been in place for a long time. SAM emphasizes that the responsibility for workplace safety rests with the employer and that the employer should work continuously with mapping the workplace risks as well as making the necessary arrangements for preventing accidents and health problems. These tasks are supported by a comprehensive set of rules and regulations as well as an organization of internal safety representatives. The authority responsible for supervising that employers abide by the rules can also initiate legal action when necessary. As the government has launched a number of inquiries related to this area, the foundations of the Vision Zero are under construction rather than being built into the work from the beginning. A number of committees containing experts are presenting reports related to parts of the government decision. It is interesting to note that this and earlier mentioned Vision Zero policies have been inspired by the vision for road safety, but the theoretical and practical foundations have not been in place to a larger degree. One question that will be addressed in the following analysis is whether it is a problem or an opportunity to issue a Vision Zero on that basis.

Analyzing Differences and Similarities in Vision Zero Components

Despite the striking similarities in terms of problems addressed (injuries) and the way the visionary goals are formulated (zero deaths, etc.), it is obvious that the five Vision Zero policies also differ significantly with regard to the preconditions needed to actually influence the development of injuries in the desired direction within each policy area. The differences are manifested in problem framing, in the monitoring of relevant facts, plus, not least, in access to means, strategies, and governing structures.

Problem Framing

Problem framing is always an important foundation for any policy aimed to address a certain problem. The framing should be scientifically anchored, broad-minded, and problem oriented. The framing helps to clarify the nature of the problem at hand, including its spectrum of determinants and potential strategies, and thus creates trust in the theoretical possibility of prevention.

In this respect, the Vision Zero policy for **road traffic** safety can be viewed as a role model. By clarifying the key mechanisms of crash violence, its transmission to human tissues, and potential to harm if human tolerance limits are exceeded, combined with a modeling of theoretically available alternatives to prevent this transfer to human bodies, there is a growing trust in the potential to prevent the problem. The question is no longer *if* the problem can be prevented, but rather in what pace and to what costs. Related to this, emphasizing the distinction between accident and injury is an important contribution. Accidents (crashes) may continue to occur due to system imperfections but do not necessarily need to cause death or severe injury. In contrast to earlier views where accidents were seen as the phenomenon to prevent, Vision Zero points out deaths and serious injuries as the undesired target outcome. Injuries are preventable even if accidents continue to occur.

In comparison, the framing of Vision Zero for **fire** safety appears less elaborated. Fire safety as an academic discipline rests historically on knowledge on fire dynamics in buildings, extinguishing techniques, and rescuing strategies. It is presupposed that a fire becomes increasingly dangerous to humans as it escalates. Recent research, however, shows that many fatalities occur already in the initial stage of a fire before it spreads to the whole dwelling. Smoldering fires in upholstered furniture may generate imminent toxic gases with rapid medical effects, and clothing fires may cause immediate life-threatening burns. Professional learning accumulates from larger fires subjected to callouts, but what kills is usually smaller fires in fabrics and furniture, some of them not even attended by rescue services. In addition, there is an obvious social dimension related to the groups at risk of being killed and seriously injured in fires. Victims typically represent medically and/or socially very vulnerable categories – an aspect of the problem that is highly overlooked in the current profession-based learning system. How the social and medical sides of the fire problem are to be addressed is not well described. To summarize, therefore, important flaws remain in this policy area with regard to problem framing and convincing preventative alternatives.

Patient injuries occur in healthcare contexts and are generally understood and explained as negative health consequences from errors or neglects during medical care in health facilities. Patient safety is seen as an integral part of the quality of care, subjected to managerial efforts in line with general principles for quality improvements. The main responsibility rests with the care provider, while the role of societal bodies is to ensure this accountability through information, advice, and enforcement. The problem framing in patient safety therefore appears comparatively transparent and understandable. But even though the problem here is rather straightforward, other challenges affect the problem framing. For instance, the Health and Social Care Inspectorate (IVO) stated in its yearly report from 2019 that one major problem today for patient safety is that progress is made so fast, concerning, for instance, medical methods and techniques. The healthcare system as a whole does not have the capacity to make sure that the improvements are spread evenly throughout the system or that they are implemented in an appropriate and informed way (IVO 2019).

Suicide is a complex phenomenon with determinants deeply rooted across societal sectors. A suicide incident is by definition self-inflicted in order to terminate life. But suicide cases are often, both practically and scientifically, blurred with adjacent phenomena, such as fatalities without known intention, cases of self-harm without intention to kill, overdose episodes among substance abusers, and the like. Underlying modifiable societal determinants remain largely unexplored. Further, there is theoretical ambiguity among researchers regarding to what degree suicide results from reasoned decision-making or from sudden situational and overwhelming circumstances (“psychological accidents”). The same ambiguity is reflected in different views on preventative strategies, spanning from mental illness identification and treatment to environmental modifications in order to reduce access to lethal means. The problem framing on suicide appears partial as it reflects a medical view of the problem and its solution, mainly, rather than a social one.

Fatal accidents in the **workplace** are, like patient injuries, easy to define without theorizing too much. Fatalities at workplaces result from falls from heights, collapsing structures, incidents with machinery, etc. The spectrum of events is more varied when compared to road traffic, but the injurious mechanisms of uncontrolled “violence” (mechanical, thermal, chemical, etc.) to the human body are similar, as well as the spectrum of measures available to prevent transfer of this harm to human bodies. The basic principle is to minimize deaths and injuries from occupational accidents by technical and organizational measures. The main responsibility rests with the employer, while the role of societal bodies is to ensure this accountability through information, advice, and enforcement. Like patient safety, the problem framing on occupational safety seems fairly transparent and understandable.

Monitoring and Surveillance

Any phenomenon subjected to systematic change should be possible to measure with regard to frequency, distribution across relevant subcategories, and development over time. When policy makers claim that something occurs too often or too rarely,

and therefore should decrease or increase, it is already implied that relevant facts exist. If the problem is injuries and deaths, the art of providing such data in a systematic manner is usually called injury surveillance (a sub-discipline of public health surveillance), or simply "injury statistics." Surveillance is a broader term that includes the collection, processing, analysis, and feedback of relevant data to those who need to know in order to take proper actions. Surveillance is aimed to serve as a driver for change. Criteria for good surveillance systems underline issues like accuracy of case definitions and inclusion criteria, validity and reliability aspects, timeliness, as well as the quality of analysis, reporting, and utilization of data. Without access to good data, it is not really possible to say what is wrong, what needs to be done, or to evaluate interventions. The preconditions concerning each of the five Vision Zero policies differ remarkably in this respect.

The policy area of road **traffic** can be seen as a role model also related to this issue. The Swedish Transport Administration has taken the issue of injury surveillance very seriously and clarified operational definitions on fatalities as well as major injuries from road traffic. Validated data collection routines are secured, combining information from the health sector and the police. The data series go back quite far in time which means that analyses on trends can be performed. It is also possible to follow subgroups so that profiled interventions can be prioritized and evaluated. Furthermore, the Vision Zero for road safety has been in place for more than two decades and the actors within the policy area have had quite some time to coordinate and also to establish a specific structure.

In **fire** safety, the situation remains more challenging. The registration of fire fatalities now follows an updated and validated routine combining data from rescue services and the health sector according to a likewise updated case definition of fire fatalities. However, there is still no case definition of major injuries from fire and no regular data collection routine established on major injuries, in spite of the priority these injuries are given in the Vision Zero.

Monitoring **patient** injuries appears even more challenging. Definitional and operational difficulties create barriers for establishing a regular comprehensive surveillance system on patient injuries. Conditions that may have contributed to a patient injury are something that often must be judged by experts in retrospect. Reporting systems based on patient compensation claims or staff reports on managerial deviations highly underestimate the real situation. Valid estimates must be derived from patient record reviews which are time-consuming and expensive. Due to these circumstances, it is currently not possible to give a clear overview of the problem and its development over time and by subcategories.

Data on **suicide** are available from the national cause of death register. Besides confirmed cases of suicide, statistics reported based on data from the register often include cases with unknown intent as well, despite striking differences in terms of demography and fatal mechanism (drowning, suffocation, poisoning, etc.) between the two categories. Adding unknown cases to the confirmed ones inflates the numbers considerably. Data on the so-called suicide attempts, available from national inpatient statistics, include a broad spectrum of injuries from self-harming and self-destructive acts, without information on whether there was an intent to really end life.

Workplace safety, finally, represents longstanding traditions on data collection and analysis for the purpose of prevention. In Sweden, data collection is based on compensation claims to the Swedish public insurance agency, plus, in severe cases, reports directly to the Swedish Work Environment Authority. Triangulation against the national cause of death register ensures reasonable validity on deaths, while underreporting exists among nonfatal cases. Incidents in informal and illegal sectors are probably more extensively underreported.

Means, Programs, and Governing Structures

Finally, we have analyzed the questions of governing and steering structures in relation to specific measures and solutions. We find that steering and governing also presuppose access to effective means, a program clarifying what needs to be done, when and by whom, in addition to a structure on how to govern the program over time in a sustainable manner. Access to means implies that important determinants should be identified and found modifiable through well-known interventions. A program is a plan for action based on grounded assumptions on how various interventions are expected to influence the target outcomes. The program should also clarify priorities over time and an allocation of responsibilities among actors. A governing structure is needed to get the program done, including implementation, coordination, performance analysis, corrective actions, and follow-ups on accountability.

Road traffic safety is a field strongly characterized by its systems approach. Road traffic is part of the transport system, aimed to provide mobility with minimal consequences for safety, health, and the environment. The overall responsibility rests with the system designers and providers, while users are expected to follow rules, pay attention, and heed to other road users. System components include road infrastructure, users, and vehicles. Measures need to be directed toward all three of these components, but priority is given to infrastructure and vehicles in order to compensate for the most unreliable component – the users. Accessibility for broad road user categories is another argument for prioritizing technical and environmental improvements, rather than placing stricter demands on users. The governance of this policy area is delegated from the government to the Swedish Transport Administration and is executed in collaboration with other relevant actors in accordance with a negotiated program where responsibilities and commitments are allocated among actors. As Vision Zero for road safety has been decided upon by the Swedish government and parliament, there is an annual reporting mechanism back to these levels on progression and further needs, intended to maintain political anchoring and support. One specific problem related to the governing of road safety is the fluctuation of the status of road safety in relation to other transport-related issues. There is a risk that this fluctuation in prioritization has effects on long-term transformation. To succeed in bringing the number of deaths down requires coordination and cooperation among many actors. Bringing all these actors together has proven quite a challenge for the Swedish Transport Administration as the lead agency. Networks are set up but the capacity of these networks has not been fully developed.

The policy area of **fire** safety appears less matured and organized. Accountability besides the individual responsibility remains unclear, both legally and in practice. The broader systems approach, like in traffic, is yet to be elaborated and fully mandated for coordination and governance to a designated body. Currently, the policy area of fire safety falls under the jurisdiction of the Swedish Civil Contingencies Agency, an agency with very limited possibilities to influence relevant conditions outside its own restricted sector. Standards for buildings and dwellings fall under other sectors, like other fire-related issues on electrical equipment, furniture, home-based healthcare and nursing, social housing, alcohol, tobacco and drugs, etc. A program is outlined, identifying a set of determinants (“indicators”) considered important to modify, but there is no overall steering apparatus established to really implement the program across sectors.

Patient safety, however, is another example of a field characterized by a systems approach, at least in writing. System components include professionals, patients, technology, and organization. The overall responsibility rests with the caregiver (organization), while single professionals are expected to comply with standards, keep themselves updated, and report deviations from safe practice. On the other hand, there is a lack of overall monitoring systems and programs allowing for broader overviews and governance. Therefore, the systems approach and the clarity regarding responsibility are issues still largely theoretical, while a concrete management structure is yet to be established. There are several actors with clear mandates to monitor and report, such as the National Board of Health and Welfare, an organization issuing important guidelines for patient safety. The Health and Social Care Inspectorate also has a monitoring role both on a general level but also directly related to patients’ complaints. On paper we have an authority providing guidelines that caregivers should follow, and we have a monitoring authority issuing actual advice on improvements, but the system is so vast that the implementation of standards is challenged.

The policy area and Vision Zero for **suicide** shows similarity with fire safety concerning the absence of a broader systems approach and a clear lead agency capable of managing the field in the intended direction. Suicide is a comprehensive societal problem rooted in broad societal developments such as economy, health, labor market, family structures, and housing, all of them conditions out of reach for single actors to change. The National Board of Health and Welfare is appointed by the government to serve as a focal point for this area. The agency has a certain mandate over the healthcare system which means that the program in practice is narrowed down to issues possible to influence through the healthcare system, like identification and treatment of depression. This approach may yield some positive results but will not affect the deeper social determinants of the problem. The critical voices from within the healthcare system and from NGOs and voluntary organizations are continuously pointing toward this narrowing down of the system itself, having clear effects on implementation and problem framing.

Workplace safety, finally, is yet another example of an area with a well-established systems approach and with a clear division of responsibilities. The Work Environment Act (1977) assigns the main responsibility to the employer.

The employer should make sure that all equipment is safe and that employees are properly informed and educated to perform the work in a safe way. The Swedish Work Environment Authority is the lead agency expected to ensure, through information and enforcement that the employer takes on the responsibility for workplace safety in a satisfactory manner. The societal steering is thus performed indirectly by regulation, enforcement, and advice, which in practice limits the possibility to directly affect the development. Recognizing that the actors within this policy area are increasingly working on an international market entailing consequences for safety, wages, and social conditions, this also has consequences for the governing and steering structures related to workplace safety.

Discussion

The analysis of the five cases shows the difficulties and challenges of governing based on a vision and in combination with the zero approach. The Vision Zero role model within road traffic safety was developed in close relation to scientific results on, for instance, crash violence and was also influenced by other events over time in Sweden. Although Vision Zero has continued to develop within this policy area and has been subjected to constant improvement, its foundation appears more solid than the other cases. In this final section of the chapter, we will, based on the empirical findings and comparative analysis, return to three questions raised in the analytical framework:

1. Does a policy have to contain specific criteria in order to be called a Vision Zero policy, and what should we normatively ask of a Vision Zero related to reaching its end goal?
2. Do these criteria have to be in place before the adoption of the Vision Zero policy or can they be developed in a continuous transformation process?
3. In the light of its diffusion all over the world, how can we distinguish Vision Zero from other zero perspectives and why is that important?

Are There Discernable Determinants for a Vision Zero and for It Being Successful?

It is obvious that the compared Vision Zero policies differ in terms of practical feasibility and thereby also in trustworthiness with regard to their possibilities to affect the specific outcomes targeted in each policy. If a policy fails to scientifically frame the problem properly, including determinants and preventability, or fails to measure its problem's frequency and severity across relevant categories and over time, or lacks fundamental instruments for change, it appears problematic to denote it a Vision Zero policy, since there is little or no chance for the policy to fulfill its mission. Doing so may instead erode public trust in Vision Zero policies in general and eventually endanger the whole idea of Vision Zero policies. In our

view, it is the visionary image in combination with a trustworthy apparatus for systematic steering toward this vision that legitimates the term Vision Zero. This in turn, with reference to our analytical framework, rests on the model we have used for our comparative analysis, i.e., in short, the Vision Zero policies are based on wicked societal problems and these problems have to be framed properly and consistently in order for the measures and solutions to work efficiently. One crucial framing regards the system itself and particularly its actors and structures. Another key element is a system of monitoring and feedback.

In order to be implementable, a policy has to be clear regarding problem, measures, solutions, and goals, as well as monitoring and governing system. This is very much true for all policies. But using visions as policy tools require additional approaches. The very essence of a vision is its ability to inspire and to affirm important societal values for an extended period of time. To transform a vision into a workable tool requires patience, and adding a zero approach to a vision necessitates coordinated efforts. Visions thus contain both an element of inspiration and an opening for transformation toward implementation in practice. Based on the analysis of the five Vision Zero policies in Sweden, we conclude that there are problems and opportunities with governing by visions. We would like to take the discussion above a bit further by identifying a number of more specific criteria that in our view are necessary in order to work with a vision based on a zero approach in relation to wicked problems within the field of injury prevention. There has to be:

- **Scientifically** determined problems and solutions (in depth and width), including its spectrum of modifiable determinants at individual, technical/environmental, organizational, and societal levels.
- A **comprehensive** approach. For a vision to be successful, it is necessary to view the society in a holistic way. This requires knowledge of what policy measures are effective together and presupposes an analysis (often referred to as “systems analysis”) of relevant actors and incentive structures. This process often leads to broader policy programs often studied using the so-called program theory.
- A **long-term** transformation process which has to include measurements and monitoring systems, follow-up and feedback routines, program evaluation, and revision.
- A **governance** structure containing a specific system for goal setting as well as commitment, coordination, and leadership, not only from the appointed authorities but from all actors with a vested interest in solving the problem at hand. Since we are here dealing with complicated problems, not only a governing structure is required but also a governance perspective where all relevant actors work together.

If applying these criteria to our five cases as they were presented when adopted, the analysis can be summarized as following:

	Road safety	Fire safety	Patient safety	Suicide	Workplace safety
Scientific foundation	Broad	Narrow	Broad	Narrow	Broad
Comprehensive approach	Broad	Narrow	Broad	Narrow	Broad
Long-term monitoring system	In place*	Insufficient	Insufficient	Insufficient	In place*
Governance system	In place*	Missing	In place*	Missing	In place*

*“In place” here means that basic functions are in place, while operational quality and effectiveness may differ considerably

Building the Vision Zero Ship at Sea?

We have concluded that the Vision Zero for road traffic safety has a more profound foundation than the other cases in many perspectives. But the question is whether it is problematic to launch a vision without the same kind of foundation. One risk is that the vision remains only on paper and never reaches the implementation stage. On the other hand, having such an ambitious vision can inspire actors to construct methods, models, and above all identifying the system within each area, especially now when there is a role model for Vision Zero. Another problematic aspect is if the methods of the role model turn out to be less effective. The rise of deaths in the road safety statistics in recent years is a concern and adds a dimension to the discussion on having a vision as a steering and governing tool in relation to wicked societal problems. However, there is an alternative way to look at the problem with premature Vision Zero policies. They can also be perceived as challenges, revealing managerial weaknesses, and prompting actions to deal with the fundamental requirements that need to be in place for rational and systematic mitigation of adverse societal outcomes. If following the key components of working with a Vision Zero mentioned above, it should be possible to avoid an empty vision.

A Conceptual Distinction

Reviewing a policy area, intended for a Vision Zero approach, by means of our criteria applied above for policy comparisons, might facilitate the identification of such structural improvement needs. There is yet no ownership or standardization on the Vision Zero concept. But given its popularity and rapid dissemination in combination with an increasing diversity with regard to contents and applications, it might be useful to seek further clarification in order to streamline the uniqueness and theoretical relevance of the concept in contrast to parallel types of policies with similar aims and applications (for an overview of improvement principles, see Hansson 2019).

Among parallel policies and concepts, zero tolerance policies may deserve special attention. Vision Zero and zero tolerance policies are often confused, or referred to interchangeably, in the public debate. The two policies are, however, quite different. The zero tolerance concept was first introduced in crime prevention based on the idea that strict police response to minor offenses would be a way to prevent major crimes. The principles were popularized by Wilson and Kelling (1982) by launching their “broken windows” theory and claiming that indulgence to minor crimes, such as breaking windows and littering, will give way for more severe nuisance and crime (Kelling and Coles 1997). The ideas gained widespread interest and were quickly disseminated to other fields, especially drug prevention. Strict and prompt punishment of any drug involvement, even minor, was expected to deter from more serious involvements. The zero tolerance policy, as applied to crime and drug prevention, has been extensively criticized for being indiscriminate and brutal (Sharkey 2018). It is also blamed for raising barriers between the police and communities (Cox and Wade 1998). In drug prevention, the zero tolerance policy has been criticized for preventing abusers from seeking medical help in critical situations and thereby contributing to unnecessary deaths in overdoses (Tham 1998). As a reaction, the so-called harm reduction strategies are now increasingly advocated as a way to save lives. Drug users are welcomed to clinics where they can get qualified medical assistance and advice without risk of being accused of criminal behavior. The approach is intended to appear forgiving and supportive instead of intolerant and punishing. This helps to clarify the important difference between Vision Zero and zero tolerance policies. While the Vision Zero policy, as first presented in traffic safety, clearly reflects a harm reduction strategy, developed in reaction to the earlier behavior-centered strategy, the zero tolerance approach is directed toward controlling human behavior entirely, moreover by repressive means. According to the Vision Zero philosophy, environments should be designed to tolerate normal deviations in human performance by allocating responsibility to system designers as well, while the zero tolerance approach maintains strict individual responsibility and proclaims intolerance to human failure.

Cross-References

- ▶ [Vision Zero in Sweden: Streaming Through Problems, Politics, and Policies](#)
- ▶ [Vision Zero on Fire Safety](#)

References

- Belin, M., & Tillgren, P. (2012). Vision Zero. How a policy innovation is dashed by interest conflicts, but may prevail in the end. *Scandinavian Journal of Public Administration*, 16(3), 83–102.
- Belin, M., Vedung, E., & Tillgren, P. (2012). Vision Zero – A road safety policy innovation. *International Journal of Injury Control and Safety Promotion*, 19(2), 171–179.

- Christensen, T., & Læg Reid, P. (2012). Governance and administrative reforms. In D. Levi-Faur (Ed.), *The Oxford handbook of governance*. Oxford: Oxford University Press.
- Cox, S., & Wade, J. (1998). *The criminal justice network: An introduction*. New York: McGraw-Hill.
- Dignum, M., Correlje, A., Groenleer, M., & Scholten, D. (2018). Governing through visions: Evaluating the performativity of the European gas target models. *Energy Research & Social Science*, 35, 193–204.
- Ekelund, M. (1999). *Varning – Livet Kan Leda till Döden!* Stockholm: Timbro. [Warning – Life can lead to death!]
- Elvebakk, B. (2007). Vision Zero: Remaking road safety. *Mobilities*, 2(3), 425–441.
- Elvebakk, B. (2015). Paternalism and acceptability in road safety work. *Safety Science*, 79, 298–304.
- Elvik, R. (2003). How would setting policy priorities according to cost-benefit analyses affect the provision of road safety? *Accident Analysis and Prevention*, 35(4), 557–570.
- Gioia, D. A., Nag, R., & Och Corley, K. G. (2012). Visionary ambiguity and strategic change: The virtue of vagueness in launching major organizational change. *Journal of Management Inquiry*, 21(4), 364–375.
- Haddon, W. (1968). Changing approach to epidemiology prevention and amelioration of trauma – Transition to approaches etiologically rather than descriptively based. *American Journal of Public Health and the Nation's Health*, 58(8), 1431–1438.
- Hallström, T. K., & Grafström, M. (2016). *Visioner som demokratiskt förändringsverktyg?* Handelshögskolan i Stockholm Stockholms centrum för forskning om offentlig sektor (Score). Decode 2016:4.
- Hansson, S. O. (2019). Improvement principles. *Journal of Safety Research*, 69, 33–41.
- Hedlund, G., & Montin, S. (2009). *Governance på svenska*. Stockholm: Santérus.
- Hood, C. (1991). A public management for all seasons? *Public Administration*, 69(1), 3–19. <https://doi.org/10.1111/j.1467-9299.1991.tb00779.x>.
- Hysing, E. (2019). Responsibilization – The case of road safety governance. *Regulation and Governance*, 1–14. <https://doi.org/10.1111/rego.12288>
- Kelling, G. & Coles, C. (1997). *Fixing broken windows: Restoring order and reducing crime in our communities*. New York: Simon and Schuster.
- Kjellén, U. (2000). *Prevention of accidents through experience feedback* (1st ed.). New York: Taylor and Francis.
- Kornberger. (2013). Disciplining the future: On studying the politics of strategy. *Scandinavian Journal of Management*, 29(1), 104–107.
- Kristianssen, A., Andersson, R., Belin, M., & Nilsen, P. (2018). Swedish Vision Zero policies for safety: A comparative policy content analysis. *Safety Science*, 103, 260–269.
- Læg Reid, P. (2011). New public management. In: B. i Badie, D. Berg-Schlosser, & L. Morlino. (red). *International encyclopedia of political science*. Thousand Oaks: SAGE Publications, Inc., pp. 1700–1704, hämtad 4 November 2018. <https://doi.org/10.4135/9781412959636.n391>.
- Lind, G., & Schmidt, K. (1999). *Leder Nollvisionen till det Trafiksäkra Samhället?* KFB: Stockholm. [Does Vision Zero lead to a traffic safe society?]
- Nihlén Fahlquist, J. (2006). Responsibility ascriptions and Vision Zero. *Accident Prevention and Analysis*, 38(6), 1113–1118.
- Pollitt, C. (2003). Joined-up government: A survey. *Political Studies Review*, 1(1), 34–49.
- Rittel, H. W. J., & Webber, M. M. (1973). Dilemmas in a general theory of planning. *Policy Sciences*, 4, 155–169.
- Rosencrantz, H., Edvardsson, K., & Hansson, S. (2006). Vision Zero – Is it irrational? *Transportation Research Part A*, 41(6), 559–567.
- Sharkey, P. (2018). *Uneasy peace: The great crime decline, the renewal of city life, and the next war on violence* (1st ed.). New York: W. W. Norton & Company.
- Sørensen, E. (2016). Enhancing policy innovation by redesigning representative democracy. *Policy & Politics*, 44(2), 155–170.

- Swedish Civil Contingencies agency (MSB). (2010). *En nationell strategi för att stärka brandskyddet genom stöd till enskilda. Redovisning av uppdrag. (Fö2009/2196/SSK)*. Stockholm: MSB. Dnr. 2009–14343. [A national strategy to strengthen fire safety by supporting individuals.]
- Swedish Government. (1997). *Nollvisionen och det trafiksäkra samhället*. Swedish Government Bill 1996/97:137. [Vision Zero and the road traffic safety society].
- Swedish Government. (2002). *Reformerad räddningstjänstlagstiftning*. Government Bill 2002/03:119. [A reformed emergency service].
- Swedish Government. (2007). *En förnyad folkhälsopolitik*. Government Bill 2007/08:110. [A renewed public health policy].
- Swedish Government. (2009). *Patientsäkerhet och tillsyn*. Government Bill 2009/10:210. [Patient safety and supervision].
- Swedish Government. (2016a). *Arbetsmiljöregler för ett modernt arbetsliv*. Committee directive 2016:1. [Work environment regulations for a modern working life].
- Swedish Government. (2016b). *Nationellt centrum för kunskap om och utvärdering av arbetsmiljö*. Committee directive 2016:2. [National center for knowledge of and evaluation of work environment].
- Swedish Government. (2016c). *En arbetsmiljöstrategi för det moderna arbetslivet 2016–2020*. Communication. Skr.2015/16:80. [A work environment strategy for the modern working life 2016–2020].
- Swedish Governments Official Investigations. (2008). *Patientsäkerhet. Vad har gjorts? Vad behöver göras? Betänkande av Patientsäkerhetsutredningen*. SOU 2008:117. [Patient safety. What has been done? What needs to be done?]
- Swedish Law on accident prevention (LSO)*. (2003:778).
- Swedish Law on Patient safety*. (2010:659).
- Swedish Parliament. (1997a). Swedish parliament protocol 1997/98:13. *Decision on Vision Zero for road traffic safety*.
- Swedish Parliament. (1997b). *Parliament transport and communication committee report 1997/98: TU4*.
- Swedish Parliament. (2014). *Arbetsmiljö*. 2013/14:AU5. Report Parliament Labor market committee. [Report *Work Environment*].
- Tham, H. (1998). Swedish drug policy: A successful model? *European Journal on Criminal Policy and Research*, 6(3), 395–414. <https://doi.org/10.1023/A:1008699414325>.
- The National Board of Health and Welfare. (2013). *Förslag till nationell strategi för ökad patientsäkerhet (A proposal for a national strategy for increased patient safety)*. Stockholm: National Board of Health and Welfare.
- The National Board of Health and Welfare and The Public Health Agency of Sweden. (2006). *Förslag till nationellt program för suicidprevention – befolkningsinriktade och individinriktade strategier och åtgärdsförslag*. [A proposal for a national program for suicide prevention]. Stockholm: National Board of Health and Welfare.
- The Swedish Health and Social Care Inspectorate. (2019). *Årsredovisning (Annual report)*. Stockholm: Directorate General, The Swedish Health and Social Care Inspectorate.
- The Swedish Work Environment Authority. (2001). *Arbetsmiljöverkets föreskrifter för systematiskt arbetsmiljöarbete*. AFS, 2001:1. [The Swedish Work Environment Authority regulations for systematic safety work in the workplace].
- The Swedish Work Environment Authority. (2016). *Kvinnors och mäns arbetsvillkor – betydelsen av organisatoriska faktorer och psykosocial arbetsmiljö för arbets- och hälsorelaterade utfall*. Kunskapssammanställning 2016:2. Författare Magnus Sverke, Helena Falkenberg, Göran Kecklund, Linda Magnusson och Petra Lindfors. Stockholm: Arbetsmiljöverket. [Women and men's working conditions – the role of organizational factors and psychosocial work environment for work- and health-related outcomes].
- The Work Environment Act. (1977). SFS 1977:1160.

- Tingvall, C., & Haworth, N. (1999). *Vision Zero – An ethical approach to safety and mobility*. Paper presented to the 6th ITE International Conference Road Safety & Traffic Enforcement: Beyond 2000, Melbourne, 6–7 September 1999.
- Tryssel, K. (2018). Läkare kritiska till nollvision för självmord. *Läkartidningen*, 115, E63W. [*The Medical Journal: Physicians are critical towards the Vision Zero for suicide*].
- Weiss, J. A., & Piderit, S. K. (1999). The value of mission statements in public agencies. *Journal of Public Administration Research and Theory*, 9(2), 193–223.
- Wilson, J. Q., Kelling, G. L. (1982). Broken windows: The police and neighborhood safety. *The Atlantic Monthly*, 22–38.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.



Responsibility in Road Traffic

5

Sven Ove Hansson

Contents

Introduction	178
What Is Responsibility?	178
Causality	183
The Multiplicity of Causal Factors	184
The Insufficiency of Cause-Effect Relationships	185
Agent Causality	186
Causes with Moral Foundations	187
The Politics of Causality	190
Responsibility in Road Traffic	192
The Traditional Approach	192
Vision Zero	195
Self-Driving Cars	196
Institutional and Professional Responsibility	198
Cross-References	200
References	200

Abstract

Vision Zero requires a new approach to the responsibility for safety. This chapter provides conceptual tools for the description and analysis of this and other responsibility issues. Distinctions between different types of responsibility are introduced, with a particular emphasis on the distinction between blame responsibility and task responsibility. The complex relationship between responsibility and causality is also delineated. This is followed by an analysis of the changes in responsibility assignments that are necessary to implement Visio Zero.

S. O. Hansson (✉)

Division of Philosophy, Royal Institute of Technology (KTH), Stockholm, Sweden

Department of Learning, Informatics, Management and Ethics, Karolinska Institutet, Stockholm, Sweden

e-mail: soh@kth.se

© The Author(s) 2023

K. Edvardsson Björnberg et al. (eds.), *The Vision Zero Handbook*,
https://doi.org/10.1007/978-3-030-76505-7_5

177

Keywords

Blame responsibility · Causality · Responsibility · Task responsibility · Traffic safety · Vision Zero

Introduction

She ran red lights at high speed and crashed into the helpless cyclist. There can be no doubt that she is responsible for the accident.

Failing brakes are responsible for numerous accidents on icy roads.

It is true that the crash was caused by the pedestrian's erratic behavior, which forced several drivers to dangerous maneuvers. But he has a severe mental disorder and is not really responsible for what he did.

In the last few years, two children have been killed in traffic accidents on their way to this school. The traffic conditions are clearly unacceptable. Something must be done. Who is responsible?

Traffic safety is one of the many social areas in which assignments of responsibility are important and often contested. They have become an even more important topic through Vision Zero, which distributes responsibilities in new ways. But as illustrated in the four examples above, we use the terms “responsible” and “responsibility” in several meanings. This chapter will begin by systematizing the major meanings of the terms. After that we will investigate the complex relationships between responsibility and causality and finally show the bearing that all this has on traffic safety and Vision Zero.

What Is Responsibility?

The most influential classification and clarification of the different meanings of “responsibility” is due to the British legal philosopher H. L. A. Hart (1907–1992). His work is therefore the best point of departure for an analysis of the concept. He identified four major meanings of “responsibility,” as the word is used in moral and legal contexts:

- By *liability-responsibility*, he meant, in a legal context, liability for punishment or for paying compensation (Hart 2008, pp. 222, 225). In most cases, liability-responsibility pertains to a person's own actions and their consequences, but there are also cases in which a person is “responsible vicariously or otherwise for harmful outcomes which he had not caused” (ibid., p. 224). In a moral context, liability-responsibility usually means that the person deserves blame, rather than punishment, but in some cases the person is “morally bound to make amends or pay compensation” (ibid., p. 225).

- By *role-responsibility*, Hart meant the “specific duties” a person obtains through occupying “a distinctive place or office in a social organization” (ibid., p. 212). His usage of “role” covers not only professional and official functions but also private roles such as those of a spouse, parent, or host.
- With *causal responsibility*, he referred to cases in which the phrase “is responsible for” can be replaced by “caused” or “produced,” without a change in meaning. According to Hart, causal responsibility can be attributed not only to human beings “but also to their actions or omissions, and things, conditions, and events.” One of his examples is: “The icy condition of the road was responsible for the accident” (ibid., p. 214).
- By *capacity-responsibility*, he meant capacity to understand, reason, and control one’s own actions. This is what we refer to with the phrase “he is responsible for his actions.” In order to have capacity-responsibility, the person must have “certain normal capacities,” namely, “those of understanding, reasoning, and control of conduct: the ability to understand what conduct legal rules or morality require, to deliberate and reach decisions concerning these requirements, and to conform to decisions when made” (ibid., p. 227).

There are close connections between moral and legal concepts of responsibility, and one may see legal responsibility as a codification of such moral responsibilities that we have agreed to impose upon each other with the force of law. Here, we will focus on moral responsibilities. Let us consider, in turn, each of Hart’s four types of responsibility.

As already indicated, Hart was aware that *liability-responsibility* may not be an ideal term in moral discussions. Whereas this form of responsibility is usually strongly connected with liability in legal contexts, in a moral context it is more closely connected with blameworthiness. Therefore, moral philosophers writing about responsibility usually prefer the term “blame responsibility” (Goodin 1987, p. 167). This terminology will be used here as well. However, it should be noted that in addition to deserving blame, a blame-responsible person may also be morally required to compensate negatively affected persons as well as to perform other acts of expiation (Hansson and Peterson 2001).

When we talk about a person as being responsible for something that she has done, we usually focus on the negative consequences of her actions. However, one can also be responsible for laudable acts. The Oxford English Dictionary has a value-neutral definition of the term as “[t]he state or fact of being the cause or originator of something; the credit or blame for something.” The European Transport Training Association hands out a yearly Safety Award, which “recognizes those responsible for excellent products or services aimed at improving road safety in the European road transport and logistics industry” (Anon 2013). This usage of the word “responsible” can be termed “praise responsibility.” Perhaps it should have a larger social role than what it has – for instance in traffic safety – but for our current purposes it can be left out of the discussion.

Hart’s “role responsibility” refers to what one has to do or achieve. Several authors have noted that his terminology tends to obscure the generality of this notion

(Cane 2002, p. 32). The word “role” suits well for legally binding responsibilities, such as those that follow with an employment contract, a marriage, parenthood, or board membership in voluntary organizations. However, it does not suit well for more informal undertakings that are usually considered to confer some responsibility, such as agreements to babysit, water someone’s flowers, walk their dog, or feed their aquarium fishes. Some authors have kept Hart’s term “role responsibility” but interpret it very widely (Dworkin 1981, p. 29). Others use the term “task responsibility,” which has a wider general meaning and obviously covers “duties, jobs or (generically) tasks,” including those that originate in informal undertakings and agreements rather than legally binding stipulations or contracts (Goodin 1987, p. 168; Cane 2002, p. 32). Here, “task responsibility” will be used as a general term for responsibility to do or achieve something.

According to Hart, most adults are considered to have *capacity-responsibility*, but it is “lacking where there is mental disorder or immaturity” (Hart 2008, p. 218). We can speak of a person as being responsible for her actions, in this sense, even if we do not know of any particular action that she is responsible for. Therefore, capacity-responsibility should be seen as an ability to be responsible, rather than as a form of responsibility per se. The legal notion of capacity-responsibility is related to the notion in medical ethics of “capacity for autonomous choice” (often also called “decision-making capacity” or “competence”), which marks the limit between those who can respectively cannot give informed consent to a medical intervention (Parker 2001; Stirrat and Gill 2005; Michaud et al. 2015). The term *capacity* (or *capacity to be responsible*) can be used for this notion. Here, we will have relatively little use for it, since issues of capacity (or ability to take responsibility) seldom arise in discussions of traffic safety. Drivers are required to have a driver’s license, which is normally only issued to adults with the requisite abilities. The protection of pedestrians lacking in the relevant mental capacities, such as children and people with mental disabilities, is an important issue in traffic safety, but it is usually discussed in terms of the risks they are exposed to rather than their capacity to take responsibility.

The way Hart uses the term *causal responsibility*, it is not entirely clear why he did not instead use the term “causality” for this notion. An icy road can certainly be the cause of an accident, and someone who should have sanded it can then be responsible for the accident, but what is gained by saying that the road itself is responsible? It would seem more clear to reserve the term “responsibility” for agents who can reason and argue and use the term “causality” for inanimate objects.

However, this becomes somewhat more complex in cases when causality is ascribed to humans or to their actions. Consider the following two examples:

Case 1:

Adam was terribly drunk and fell asleep on the kitchen floor. Two of his friends moved him to the floor of an adjacent room, just to keep him out of the way. Susan entered the room without noticing him, stumbled over his legs, fell on a chest of drawers, and broke her nose.

Case 2:

Adam was lying on the floor when Susan entered the room. He stretched out his leg to tease her. She did not notice, stumbled over his leg, fell on a chest of drawers, and broke her nose.

In the first case, Adam caused Susan’s broken nose in the same way as it could have been caused by a sack of potatoes, or some other inanimate object. In the second case, his causal role was different, since he made a decision – namely, to stretch out his leg – that had a crucial causal role. This is a role that only an agent can have, and “his, her or its agency serves to explain” the pertinent outcomes, which “can therefore plausibly be treated as part of the agency’s impact on the world” (Honoré and Gardner 2010). It is not unreasonable to use the term “causal responsibility” in this case (contrary to the case with the icy road, mentioned above). However, the term “agent causality” will be used here instead. The reason for this terminological choice is that agent causality does not necessarily imply responsibility in any moral or legal sense. If Alyona causes Diego’s death while doing her very best to save him from a life-threatening danger, then she is an agent-cause of his death, but not necessarily morally or legally responsible for it.

As summarized in Table 1, we have renamed and adjusted Hart’s (mainly legal) terminology to make it more suitable for moral investigations. Notably, only two forms of responsibility remain, namely, blame and task responsibility. We have assigned other names to Hart’s other two responsibility concepts, names that do not designate them as forms of responsibility.

This reduction to two forms of responsibility is by no means original; to the contrary it is a common approach in the literature on responsibility. (A notable exception is Gerald Dworkin, who listed three major types of responsibility in an influential article: role responsibility, causal responsibility, and liability responsibility (Dworkin 1981).) However, it is common to use other terms for blame and task responsibility, namely, terms that indicate temporal relationships. Blame responsibility is often called “backwards-looking responsibility,” “retrospective responsibility,” or “historic responsibility,” whereas task responsibility is referred to as “forwards-looking responsibility” or “prospective responsibility” (van de Poel 2011; Duff 1998; Cane 2002, p. 31). Unfortunately, this temporal terminology is somewhat misleading. We can refer in retrospect (“historically”) to the task responsibility of medieval physicians to treat patients during an epidemic in spite the grave risks to themselves (Huber and Wynia 2004). Then we have a backwards-looking perspective on a (previous) task responsibility. We can also consider

Table 1 A comparison of terminologies for responsibility-related concepts

<i>Hart’s terminology</i>	<i>Our terminology</i>
Liability-responsibility	Blame responsibility
Role responsibility	Task responsibility
Capacity-responsibility	Capacity (to be responsible)
Causal responsibility	Agent causality

prospectively whether our actions and omissions will in the future give rise to blame responsibility (Hansson 2007). Then we have a forwards-looking perspective on (future) blame responsibilities. Strangely, the latter but not the former case is called “historical responsibility” in Cane’s (2002, p. 31) terminology. The “blame” and “task” terminology does not run into these difficulties, and it will therefore be used here.

Blame and task responsibility are often closely connected to each other. One type of connection between them ensues when a failure to fulfill a task responsibility gives rise to a blame responsibility. If I have promised to water your garden while you are away, then I have a task responsibility to do so. If I fail to do it, then I am blame responsible for this failure. Another type of connection arises when a wrongful action gives rise both to blame responsibility and to a task responsibility to improve one’s future behavior. If I disturbed my neighbor’s sleep by playing music too loud in the night, then I am not only blame responsible for the disturbance but also task responsible for not repeating it in the future.

However, in more complex social situations, blame and task responsibility do not always follow each other that closely. For instance, suppose that a speeding motorist runs over a child crossing a road on its way to school. In the subsequent trial, the driver will be held (blame) responsible for the act. And of course the driver is (task) responsible for not driving dangerously again. But that is not enough. We also need to prevent the same type of accident from happening again, with other drivers. This is not something that the culpable driver can do. Instead, measures are needed in the traffic system. We may have reasons to introduce traffic lights, speed bumps, or perhaps a pedestrian underpass. The task responsibility for these measures falls to decision-makers such as public authorities. In cases like this, blame and task responsibility part company.

It has sometimes been assumed that the assignment of blame responsibility is some sort of zero-sum game, so that more responsibility for one party must always be linked to less responsibility for someone else. This has often taken the form of a principle of “proportionality,” according to which “[a]n agent’s moral responsibility for an outcome is proportionate to her actual causal contribution to the outcome” (Bernstein 2017). There are of course cases in which actions that make one agent more blame responsible also reduce the blame responsibility of some other agent(s). However, this does not hold in general. This can perhaps be most clearly seen from cases of overdetermination. If two persons simultaneously shoot a non-threatening victim, and each of them delivers a deadly shot, then this certainly does not mean that each of them is only half as blame responsible as if the other had not pulled the trigger (Bernstein 2017; Moore 1999, p. 10). Similarly, if two motorists drive into a four-way crossing at the same high speed, causing a crash that would also have occurred if only one of them had driven too fast, then neither of them is relieved of his blame responsibility by the other’s wrongdoing. Thus, blame responsibility is not a zero-sum game.

A parallel argument applies to task responsibility. There are cases when task responsibilities can be transferred from one person to another. This typically

happens when people share a task. For instance, if Erol and Haluk take turns helping their old mother on alternate weeks, then as long as this arrangement lasts, each of them has arguably only half the task responsibility that he would have had if his brother did not help. However, there are other cases in which one person's responsibility does not decrease the responsibility of others. If government takes more responsibility for reducing traffic accidents, for instance by making roads safer, then this does not reduce the responsibility of individual drivers to drive safely. (Instead, it makes it easier for them to fulfill that responsibility.) Although both blame and task responsibility can sometimes be shared, neither of them is in general "like a pie that is to be divided between people that each will have a smaller or larger share" (Verweij and Dawson 2019, p. 100). Our responsibilities are influenced by what other people do and undertake, but often in much more complex ways than that.

Causality

Both task responsibility and blame responsibility are closely connected with causality. Task responsibilities are normally assigned to persons who are presumed to be able to fulfill the task in question successfully. People are held blame responsible for their actions and for outcomes that they have caused or at least causally contributed to (Shaver 1985; Cane 2002; Moore 2009; Mumford 2013; Bernstein 2017, p. 165). However, there are exceptions to this. The law has "pockets of strict liability," by which is meant liability that can be assigned even without any causal contribution (Moore 2009, p. 21n). In many legal systems, owners of a dangerous animal are held blame responsible for injuries inflicted by the animal, and companies are held responsible for defective products, regardless of fault or causality. In a somewhat analogous manner, government ministers and leaders of public companies and other large organizations are often held (morally) blame responsible for wrongdoings by employees.

In moral and legal philosophy, it is often assumed that causality is a well-defined and value-independent phenomenon that can serve as a suitable fact base for value-laden concepts such as responsibility. However, this is a gross oversimplification that does not take into account the complexities of our concept of causality.

The usual approach to causality, applied in moral philosophy as well as in everyday life, takes causality to be constituted of (binary) cause-effect relationships. Such relationships are useful for describing many of the events that we observe around us. For instance, Carina throws a ball at the window, and the window breaks. This is a relationship between two events, a cause and an effect. Her act of throwing the ball is the cause, and the breaking of the window is the effect. In a simple, causally determined world, everything that happened would be the outcome of such cause-effect relationships. But that is not the type of world we are living in. The actual workings of the physical universe deviate from that description in at least two important ways.

The Multiplicity of Causal Factors

The first of these deviations is that instead of a single cause, there are usually several causal factors contributing to an effect. For instance, suppose that Nadja won a game of chess against Boris. We can treat her victory as an effect. What was the cause of this effect? In fact, all of the following can – all at the same time – be reasonable answers to that question:

It was because of her brilliant queen sacrifice in move 22.

It was because she has carefully studied Rudolf Spielmann's book, *The Art of Sacrifice in Chess*.

It was because Boris made a mistake in move 21 that opened up several winning strategies for her.

It was because Boris had a migraine and did not play at his best.

...

This is how it usually is. As was pointed out by John Stuart Mill ([1843] 1996, pp. 327–334), there are normally several causal factors that contribute to the production of an effect. But as he also pointed out, we seldom try to deal with them all on an equal basis. Instead, we tend to select only one of them and call it “the cause.” It is not uncommon that different persons choose different causal factors as “the cause.” In this case, it would be no surprise if Nadja sees her studies of Spielmann's book as “the cause” of her victory, whereas Boris considers his migraine to be the true cause. If the game is published in a chess magazine, readers can be expected to see either move 21 or move 22 as the cause of her victory.

Our choice of “the cause” among the causal factors that (jointly) lead up to an event depends on our perspective on that event and its antecedents. There is usually no single perspective that is more “right” than the others, and therefore there is no “right answer” to the question what “the cause” of an event is. This can also be seen from a classic example, namely, the cause of cholera. If you ask a bacteriologist what causes that disease, you will probably be told that it is caused by the bacterium *Vibrio cholerae*. If you ask an epidemiologist the same question, you will learn that it is caused by lack of proper sanitation (Rizzi and Pedersen 1992). They are of course both right. Their answers do not reveal a difference in opinion; they just put emphasis on different components in a complex causal process. The two answers can and arguably should coexist since they are useful in different contexts. A physician treating a patient with cholera has reasons to focus on the microbiological cause, whereas the cause mentioned by the epidemiologist should be at focus in preventive work. Attempts to make one of these two causal factors “the cause” for all purposes will render us less capable to solve urgent practical problems. (“Don't worry about sanitation. Cholera is caused by *Vibrio cholera*, nothing else.”)

Cause selection is a rather complex process that does not seem to be governed by a single rule. It can be likened to the use of concentrated lighting on a theater stage. With a spotlight, all the light can be put on a small part of the stage. Often, there are several artistically reasonable ways to do this, representing different perspectives on the unfolding drama. Similarly, cause selection can be performed in many different ways.

We sometimes single out one among all the causes of some event and call it “the” cause, as if there were no others. Or we single out a few as the “causes,” calling the rest mere “causal factors” or “causal conditions.” Or we speak of the “decisive” or “real” or “principal” cause. We may select the abnormal or extraordinary causes, or those under human control, or those we deem good or bad, or just those we want to talk about. I have nothing to say about these principles of invidious discrimination. (Lewis 1973, pp. 558–559)

The indeterminateness and lack of objective grounds for our choice of “the” cause among the causal factors has often been referred to as a “context sensitivity” of causal claims (Tarnovanu 2015, p. 68). However, it seems to be less a matter of the context than of perspectives and expectations, which may differ within one and the same context. The crucial conclusion we can draw from the multiplicity of causal factors is that our assignments of cause-effect relations depend not only on objective factors in the world but also on our perspectives and expectations.

The Insufficiency of Cause-Effect Relationships

As already indicated, the standard approach to causality, which is based on binary cause-effect relationships, also has another, even more serious problem. The problem is that cause-effect relationships only provide us with an incomplete picture of the world. Obviously, many of the interconnections that hold between different events at different points in time can be adequately accounted for with the cause-effect pattern. However, there are also important interconnections that do not fit into this pattern. In the context of natural science, this was pointed out by Bertrand Russell, who observed that “oddly enough, in advanced sciences such as gravitational astronomy, the word ‘cause’ never occurs” (Russell 1913, p. 1).

In the motions of mutually gravitating bodies, there is nothing that can be called a cause, and nothing that can be called an effect; there is merely a formula. Certain differential equations can be found, which hold at every instant for every particle of the system, and which, given the configuration and velocities at one instant, or the configurations at two instants, render the configuration at any other earlier or later instant theoretically calculable. (Russell 1913, p. 14)

Notably, the differential equations that Russell referred to have a central role both in Newtonian mechanics and in the relativity theory that replaced it. In pre-Newtonian mechanics, cause-effect relations were sufficient. This is exemplified by the clock-work universe of René Descartes, in which nature operated in the same way as “the movements of a clock or other automaton follow from the arrangement of its

counter-weights and wheels” (Descartes [1632] 1987, p. 873). In Newtonian physics, in contrast, movements emerge from complex interactions between a large number of bodies, all of which influence each other simultaneously.

Modern physics relies even more on mechanisms not describable in terms of binary cause-effect relations than the physics that Russell referred to (Kuhn 1971; Hausman and Woodward 1999). Furthermore, social science has followed physics in adopting models in which the flow of events is determined by simultaneous mutual influences that cannot be adequately described in terms of the stepwise production of effects in a causal chain. This applies for instance to equilibria in economics. Similar complex interactions are also discussed in other areas of social science, such as political and organizational science, although usually not in terms of equation systems (Dent 2003). An account of complex social phenomena that is restricted to binary cause-effect relationships will lack much of the explanatory power of modern social science (Berger 1998, p. 324).

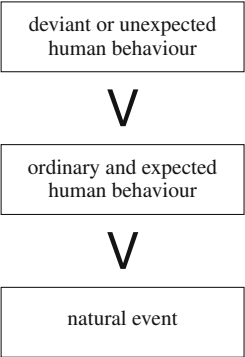
In summary, binary cause-effect relations are not sufficient to describe all the interconnections that there are among objects and events in the world. This makes it necessary to distinguish between two notions of causality. One of them is constituted by binary cause-effect relationships. We can call it *CE-causality*. The other, more general notion of causality refers to the combination of all the various types of interdependencies that obtain among objects and events in the world. We can specify it as consisting of all the connections between different events in the world by which what happens in some points in space-time restrict or partly determine what happens at other points in space-time. We can call this patterns-of-determination causality or *PoD-causality*. It is a feature of the world we are living in. CE-causality is one of the means by which we try to describe it. Newtonian and relativistic mechanics are other such means. Thus, CE-causality is a model of PoD-causality and in fact a rather crude and incomplete model. As always, it is important to distinguish between the real world and the models of it that we have created.

In this section we have made two important observations about CE-causality, which is our common notion of causality: (1) There is no objective ground for our selection of “the cause” of an event among a bundle of causal factors that contribute to it, and (2) binary cause-effect relations are insufficient to account for all the interconnections that prevail among objects and events. In combination, these two insights should be sufficient to caution us against the common assumption that responsibility can be founded on an objective fact base, consisting of cause-effect relationships. Additional reasons to be wary of that assumption will emerge from studying agent causality, the special type of causality that is particularly relevant for ascriptions of responsibility.

Agent Causality

One of the factors that affect our choice of “the cause” among causal factors is a tendency to focus on human actions whenever such actions have contributed to the effect. Therefore, agent causality has a certain priority in our account of causes.

Fig. 1 General tendencies in our choice of “the cause” of an event, among the causal factors that contribute to it. Deviant and unexpected human behavior has the largest chance to be designated “the cause.” Ordinary and expected human behavior comes second, followed by natural events



Where there is only one human causal factor, no matter how small, that factor is potentially significant (and by human causal factor here I am referring to both human action and inaction). If it is also relatively close along the chain of causation that leads to the injury, we are very likely to assign causal responsibility to it. Indeed, even a very small human causal factor may bear causal and therefore moral responsibility if there is no other human causal factor to bear it. (Reiff 2015, p. 393)

Furthermore, if there is a choice among several causal factors exhibiting human actions, then we tend to give priority to actions that stand out as in some respect deviant or unexpected. This should be clear from the following example:

Lora is driving on a country road in Kent, keeping to the left side of the road. Rose is driving in the opposite direction on the same road, but keeping to the right side. They meet at high speed in a curve and are both severely hurt.

Most of us would say, without doubt, that Rose caused the accident, since she drove on the right side of the road in left-hand traffic. But suppose that exactly the same course of events had taken place on a road in Hauts-de-France, on the other side of the Channel. Then we would have held Lora to have caused the accident, since she drove on the left side in right-hand traffic. Hence, we tend to consider deviant and unexpected behavior, rather than ordinary and “normal” conduct, as “the cause” of an event. These priorities are summarized in Fig. 1.

Causes with Moral Foundations

One of the ways in which human behavior can be deviant, or diverge from our expectations, is by departing from our moral norms. Perhaps surprisingly, moral aberrations often have a crucial role in determining our choice of “the cause” among a set of causal factors. There is a considerable amount of psychological research showing the role of norms in causality ascriptions (Willemsen and Kirfel 2019). For our present purposes, it is probably more useful to show this with the help of a couple of illustrative examples.

Due to massive rainfalls, a segment of the river bank has been undermined, and anyone entering the area runs the risk of being drawn into the dangerous rapids.

Case 1: Charles is well aware that a large part of the foundations of the river bank has been swept away. In spite of this, he recommends Andrew to go all the way down to the river to look for fish. The bank collapses, and Andrew drowns in the rapids.

Case 2: Charles has no means of knowing that the river bank is damaged. He recommends Andrew to go all the way down to the river to look for fish. The bank collapses, and Andrew drowns in the rapids. (Hansson 2022)

In the first case, it seems reasonable to claim that Charles's ill-considered advice was "the cause" of Andrew's death. We would probably not hesitate to say that Charles "caused" the accident. In the second case, such a statement would appear much more problematic. Although Charles's advice is a causal factor in both cases (presumably, the accident would not have happened without it), we are much more hesitant to call it "the cause" of the accident in the second case. The crucial difference seems to be that we consider him morally culpable in the first but not in the second case.

Despite her parents' advice to the contrary, Anne goes for a long walk on an unusually cold winter day, wearing only thin summer clothes and no coat or jacket. Three hours later she calls her parents from a hospital, where she is treated for severely frostbitten toes. "It's so unfair", she sobs. "Why should this happen to me of all the thousands of people who were out there in the streets?"

"But dear Anne", says her mother. "I am sure they all had much warmer clothes than you. In this weather it is almost certain that you will have a cold injury if you dress like you did. No doubt, your way of dressing was the cause of your injury." (Hansson 2022)

What Anne's mother says makes sense and would be fairly uncontroversial. In this case, there are two obvious causal factors: the cold weather and Anne's decision to take a long walk in thin clothes. The first of these is a natural event, whereas the second is a consciously chosen human activity. As indicated above, we have a strong tendency to prefer human actions to natural events as "the cause" of something that happens.

Despite her parents' advice to the contrary, Anne goes for a long walk in the late summer evening, wearing an unusually skimpy dress. Three hours later she enters a police station, weeping incessantly, to report a rape.

In the trial three month's later, the defendant's attorney says: "There were thousands of women out in the streets that evening. In all probability, Anne was the only one who wore such an unusually revealing dress. We have just heard my client telling us that this is what made him approach her – admittedly in a somewhat pushy manner – rather than someone else. Given what we know about young men in this city I am convinced that if he had not approached her in this manner, then someone else would have done so. It is therefore obvious that her dress was the dominant causal factor that led up to the interactions that we are here to clarify. I do not hesitate to say that her way of dressing was the cause of what happened." (Hansson 2022)

This example is in some respects similar to the previous one, but there is a crucial difference: the most obvious alternative to characterizing Anne's behavior as "the cause" is to assign that role to the actions of the rapist. The reason why we find the attorney's causal claim to be preposterous is that the rapist's actions are by all standards

incomparably more morally reprehensible than any choice of clothes that a person could make. Again, our moral appraisal determines our choice of “the cause” of what happens.

Let us now consider a couple of examples from road traffic.

Case 1: A man steps out into a motorway where no pedestrian access is allowed. The driver of an approaching car tries but fails to stop, and the man is killed.

Case 2: A man walks out into the street on a pedestrian crossing. The driver of an approaching car tries but fails to stop, and the man is killed.

Even if what happens physically is exactly the same in the two cases, we are much less willing in case 2 than in case 1 to describe the pedestrian’s behavior as “the cause” of the accident. The reason for this is of course that he is morally less at fault in the latter case.

A drunk driver loses control of the car, which hits and kills a woman walking on the pavement.

In this case, we would typically describe the drunk driving as “the cause” without even considering other causal factors. But there are at least two other causal factors at play. One of them is the pedestrian’s choice of a place to walk. This accident would not have happened if she had been somewhere else at this moment. However, since her action is morally unassailable, it is a much less plausible candidate for “the cause” than that of the driver.

The other factor that we should consider in this case is a technical feature of the car, namely, that it was so constructed that an inebriated person could start and drive it (since it had no alcohol interlock). If we treat this as just a physical fact, then it cannot compete with the driver’s behavior for the position as “the cause” of the accident. However, the car is not just a physical object but also a designed product, and decisions have been made on what safety features it should and should not be provided with. If we shift our perspective from this particular accident to the large number of accidents that involve drunk driving as an essential component, then these decisions might very well be a plausible choice for “the cause” (Cf. Grill and Fahlquist 2012).

All these examples contribute to making it clear that agent causality is strongly connected with moral assessments of actions. This was observed already by Ludwig Wittgenstein, who said: “Calling something ‘the cause’ is like pointing and saying: ‘He’s to blame!’” (Wittgenstein 1976, p. 410).

For good reasons, we want to separate our discussions and deliberations on facts as far as possible from our moral beliefs (Hansson 2018). This, however, does not seem to be fully achievable since, as we have seen in the above examples, moral concerns are often decisive for what we choose to call “the cause” of an event. This is a most undesirable conclusion, since it appears to entail that we are stuck in a kind of moral morass with no means to reach a stable factual ground. It does not seem possible to systematize our moral thoughts in a precise and well-ordered manner if our factual statements about human action are inextricably coalesced with our moral assessments. It should therefore be no surprise that many philosophers have

toned down the influence of moral appraisals on causal claims or maintained (unrealistically) that our everyday concept of causality can be purged of its moral contents. (See Reiff (2015) and Tarnovanu (2015) for unusually clear statements of the issue and good selections of references.)

But there is another way out, which becomes obvious once we have realized that CE-causality is only a model for describing factual connections in the world. The problems that we encountered in the above examples were all connected with attempts to identify “the cause” in an objective way. The failure of attempts to do so does not show that we have no means of separating factual and moral assertions from each other; it only shows that the highly simplified single-cause variant of CE-causality (which is usually presumed in moral discussions) does not provide us with means to make such a separation. The chances of achieving the separation will be much better if we replace the search for single causes by attempts to identify multiple causal factors or even, when necessary, turn to models of simultaneous interactions that go beyond what can be described in terms of binary cause-effect relations (Hansson 2010). This should come as much less of a surprise to accident analysts than to moral philosophers. Accident analysis has a long tradition of searching for multiple causal factors rather than a single cause. The focus is usually on causal factors whose elimination is predicted to be feasible and to reduce the risks of future accidents. Such causal factors are usually called “root causes” (Parry 1991; Rooney and Vanden Heuvel 2004; Boyd 2015). This is arguably a somewhat misleading terminology. “Target causes” would be a better term, since the “root causes” are selected to be targeted in subsequent safety work.

The Politics of Causality

The strong connection between causality and responsibility has important consequences for political discourse and action. If the public conceives an activity as the cause of something undesirable, then chances are high that they will hold those who perform that activity responsible and require changes in their behavior. This gives rise to a “politics of causality,” i.e., attempts by different actors to influence public perceptions of causality.

The most common strategy in causality politics can be called *backgrounding*. It consists in attempts to move, in the public’s perception, a causal factor as far into the presumably unalterable background as possible. Backgrounding is usually performed on behalf of social actors whose activities give rise to a causal factor for some socially undesirable outcome. They try to turn away the public’s attention from their own contribution, often by pointing at some other causal factor for which someone else can be blamed. Organizations that contribute to health risks and other dangers are particularly active in backgrounding. Tobacco companies are a prime example. Most of their victims became addicted before reaching the age of majority. In order to disclaim responsibility for the massive lethal effects of their products, these companies claim that “the cause” that a person smokes is a free and voluntary choice by herself.

The opposite of backgrounding is *foregrounding*, the process of attracting attention to causal factors that were previously parts of the unheeded background. Foregrounding is a strategy often adopted by social critics who wish to put causal factors on the agenda that were previously treated as unalterable parts of the social structure. One important example is the changed attitude to workplace health and safety that was achieved in the late nineteenth century by trade unions and public health activists. Previously, dangerous working conditions were treated as unavoidable, and workplace accidents were blamed on the victims. It is now generally accepted that workplace accidents are caused by dangerous working conditions, which employers are responsible for eliminating. Currently, foregrounding is an important part of public health efforts aimed at risk factors such as smoking, malnutrition, and obesity. In these cases, foregrounding consists in looking beyond the choices of affected individuals and addressing the “background” conditions under which these so-called lifestyle choices are made.

To illustrate the politics of causality, let us consider one hypothetical and one actual example. The hypothetical example is as follows:

A manufacturer of chain saws sells a model with a very strong motor. The user regulates the speed of the chain by pressing a handle. If the handle is pressed to the bottom, then the chain will move so fast that the user cannot control the saw, and there are grave risks both to the user her- or himself and to people in the vicinity. The saw has an instrument on which users can see if they press the handle too hard, and it is legally prohibited to pass certain marks on that instrument. But in spite of this, accidents are common, and hundreds of people die every year due to chainsaws being run at too high speeds. (Hansson 2022)

I have presented this example in various lectures and discussions and as yet never encountered a person who claimed that these accidents were caused by careless users of the saw. We seem all to agree that these accidents should be causally attributed to the dangerous construction of the saw. This causal attribution supports the standpoint that the manufacturer is responsible for the accidents and should therefore urgently provide the saws with a speed limiter that prevents them from being run at too high speeds.

Let us now turn to the actual example:

A manufacturer of motor vehicles sells a model with a very strong motor. The user regulates the speed of the vehicle by pressing a pedal. If the pedal is pressed to the bottom, then the vehicle will move so fast that the user cannot control it, and there are grave risks both to the user her- or himself and to people in the vicinity. The vehicle has an instrument on which users can see if they press the pedal too hard, and it is legally prohibited to pass certain marks on that instrument. But in spite of this, accidents are common, and hundreds of thousands of people die every year due to motor vehicles being run at too high speeds.

In this case, we tend to consider the accidents to be caused by the users (drivers), and consequently, the consumers rather than the manufacturer are held responsible for the accidents. Therefore, as noted by Christer Hydén, “[t]he most obvious measure to treat non-compliance of speed rules – the vehicle speed limiter – is not on the agenda yet” (Hydén 2019, p. 4). Estimates based on experiments with speed limiting devices

indicate that obligatory speed limiters have the potential to reduce road fatalities by about 25–50% (ibid., p. 5). However, such a measure would not be uncontroversial. In 2019, an automobile manufacturer announced that it will block speeds above 180 km/h on all their new cars, except emergency vehicles. This is a high limit that will have no impact whatsoever on non-criminal driving, but nevertheless a motor journalist made a failed attempt to start a campaign against the decision (Nilsson 2020a, b).

Notably, there is no “objective” or mechanical difference between chain saws and motor vehicles that justifies the difference between our assignments of causality in the two cases. Instead, the contrast between our judgments in the two cases reflects our customs and conventions concerning two types of technological devices. It is the politics of causality, rather than the causal structures themselves, that differs between the two cases.

Responsibility in Road Traffic

In this final section, we are going to apply what we have found out in the previous sections about responsibility and causality to safety in road traffic.

The Traditional Approach

The approach to responsibility for road safety that prevailed throughout the twentieth century has been well described as follows:

Historically, road accidents have been treated as isolated incidents caused by bad drivers and as an unfortunate side effect of increased mobility. Consequently, responsibility has been ascribed to individual road users whose behavior government responses have sought to influence through education, regulation, and control. (Hysing 2021)

This approach may have psychological advantages. According to Elaine Walster, it can be reassuring to categorize a serious accident as the victim’s fault, since we can then “assure ourselves that we are a different kind of person from the victim, or that we would behave differently under similar circumstances, and we feel protected from catastrophe” (Walster 1966, p. 74). However that may be, this approach has a most serious disadvantage: its exclusive focus on mistakes by individual road users tends to block considerations of efficient measures that would reduce injuries and fatalities.

Well into the 1960s, it was generally accepted that traffic safety was all about *crash avoidance*. Governments, automobile manufacturers, insurance companies, and motorist organizations all agreed that it was the drivers’ responsibility to avoid all collisions. The manufacturers’ responsibility was limited to making this possible by delivering vehicles with adequate mechanisms for steering and braking that were reliable enough to make sure that the driver would not suddenly lose control.

Similarly, the responsibility of road managers was limited to providing a reasonably smooth road without undetectable obstacles. Even in cases when a crash could be linked to a mechanical failure, the blame was often put on the driver for lacking maintenance (Wetmore 2004, pp. 380–382).

In the 1960s, after considerable struggles, the crash avoidance approach was supplemented with requirements of *crashworthiness*. Since crashes were unavoidable – and rising in numbers – manufacturers were now held responsible for reducing the consequences of crashes. This led to the introduction of seat belts, crumple zones, and other life-saving technologies (Wetmore 2004, pp. 383–389). Today, after more than half a century of improvements in crashworthiness, cars are much safer than they once were.

However, the demands of crashworthiness did not lead to a shift in the ascription of responsibility for crashes and their consequences. True, the responsibilities of manufacturers were extended. They now had to deliver cars equipped not only with reliable mechanisms for steering and braking but also with equally dependable crashworthiness features such as seat belts. However, it seemed – and still seems – to be assumed that the vehicle manufacturer has satisfied all its responsibilities when it has delivered a vehicle that satisfies all the legal safety requirements. For what happens after delivery, road users are still held almost exclusively responsible, even if improved or additional safety features could have prevented deaths or injuries.

An interesting example of this can be found in an article from 1978 by two British psychologists (Howarth and Repetto-Wright 1978). They reported a pattern that they had found in official documents about accidents involving child pedestrians: Such accidents were usually considered to be caused by the child's behavior. In police reports, the most common explanation of these accidents was that the child "ran heedlessly into the road." Courts tended to conclude that "in the circumstances there was nothing the driver could do," and consequently the driver was acquitted and considered blameless (*ibid.*, p. 10). The same approach was implicitly taken by road safety experts, who advocated training of children as the most important counter-measure against these accidents.

However, the two authors had made observations of children crossing roads and found that the description of their behavior as "heedless" was far from accurate. Children were typically highly aware of the traffic, and often afraid of crossing roads, but they sometimes made mistakes such as misjudging the speed of a vehicle or not noticing a vehicle because of their close attention to another vehicle. In the moments before an accident, the situation was "surprisingly symmetrical. The child can see the danger but makes the wrong judgement: the driver can see the child but misjudges what the child will do. In these circumstances," the authors said, "it is rather odd and indeed discreditable to absolve the driver from responsibility for his misjudgement but to blame the child" (*ibid.*, p. 10). Noting "how difficult it is to change the behaviour of children on the roads" (*ibid.*, p. 11), they proposed that drivers "must be regarded as at least equally responsible for these accidents, and we must now ask what could be achieved by altering their behaviour" (*ibid.*, p. 12). It was necessary to "*redefine* the responsibility of drivers for pedestrian accidents," for

the simple reason that drivers “have the greatest power to reduce these accidents” (ibid., p. 13).

This shift of responsibility from children to adults was clearly a step forward, but interestingly, the two authors explicitly dismissed proposals to reform the traffic system in order to protect children against accidents. They noted that there were people who wanted to place “our chief reliance on engineering measures to keep pedestrians and vehicles apart,” either by constructing “controlled crossings, bridges or underpasses” or by adult accompaniment or the provision of school buses. They considered all these proposals to be unrealistically costly, but they also rejected them on more principled grounds. These proposals were, as they saw it, based on the assumption that “neither pedestrian nor driver are [sic] at fault,” a standpoint that they equated with the view that “no-one is to blame.” They expressed relief that “[f]ortunately most people in this country are not willing to take up such an extreme ideological point of view” (ibid., p. 11). Remarkably, they did not mention lowered speed limits in this context. (However, in a later article, the main author mentioned that driver education should include the advice to slow down when one sees a child wishing to cross the road; Howarth 1985, p. 176.)

The focus on the road user’s individual responsibility is still remarkably strong in the traffic safety literature. That literature is still replete with claims that the vast majority of traffic accidents, typically around 90%, are caused by human failures (Algora-Buenafé et al. 2017, p. 240; Santosa et al. 2017; Harantová et al. 2019). This claim is also prevalent in the (remarkably small) ethical literature on traffic safety. For instance, Meshi Ori writes:

It is well established in traffic safety literature that human factors are the predominant causes of traffic crashes. Obviously there are physical, and probably social and cultural aspects that count as contributing factors to the causes of traffic crashes, but those are marginal and depend on the way the driver/rider is influenced by them. (Ori 2014, p. 356)

However, as we saw above, no one can establish what the “predominant causes” of traffic accidents are, for the simple reason that the designation of some causal factors as “causes” or as “predominant” cannot be done in an objective way. In accident investigations performed under the assumption that vehicles complying with the legal regulations are beyond criticism, human failures will be the predominant causal factors. If we instead assume that human mistakes are inevitable, and investigate how the technology reacts to such mistakes, then the causal analysis will have a different outcome.

In addition, the ethical literature on traffic safety contains standpoints that go even further than the technical traffic safety literature in assigning responsibility to individual road users. In his often-quoted 2004 paper on traffic accidents, Douglas Husak observes that “personal vehicles cause tremendous amounts of harm” and adds that “much of this harm is caused culpably” (Husak 2004, p. 351). Without even considering other options, he assigns this culpability entirely to individual drivers and proceeds to discuss “*moral* questions about the use of personal motor vehicles.” In doing this, he goes beyond “the trivial observation that many motor vehicle accidents

result from speeding, alcohol impairment, or some other kind of unlawful mode of operation” (ibid., p. 351). In his view, even careful and law-abiding driving involves so large risks for other persons that driving “for frivolous purposes” is immoral (ibid., p. 362). This would include “traveling across town to patronize a new bar or restaurant,” going “from one outlet to another” to find a cheaper product, as well as all forms of “purely recreational” driving. He also finds it culpable that “[m]any persons elect to live [at] great distances from their place of employment” so that they have to travel longer than necessary to work (ibid., p. 361).

Husak himself recognizes the crucial weakness of assigning, as he does, causality and responsibility for traffic accidents exclusively to the individual road users.

I have no illusions that the general public will be receptive to my proposals. Pleas to curb driving are likely to be met with ridicule and hostility. (Husak 2004, p. 370)

He is not alone in this insight. The limitations to what can be achieved by attempts to change road users’ behavior were a major factor leading to a new approach that aims for radical improvement of the traffic system.

Vision Zero

This new approach received its first official formulation in 1997 when the Swedish Parliament adopted Vision Zero as the overarching framework for road traffic safety in the country (Rosencrantz et al. 2007; Belin et al. 2012). The basic assumption of Vision Zero is that “from an ethical standpoint, it is not acceptable that any people die or are seriously injured when utilizing the road transportation system” (Government Bill, 1996/1997:137, p. 15). All serious accidents are considered to be unacceptable, and efforts to reduce the number of fatalities and serious injuries must continue assiduously as long as accidents still occur. This cannot be achieved with the traditional approach that assigns almost the whole burden of responsibility to drivers and other road users. Therefore, Vision Zero makes the designers and implementers of the transport system responsible for eliminating human deaths and injuries. In the terminology introduced above, the movement for Vision Zero is an unusually clear example of a movement for the foregrounding of previously backgrounded causal factors.

The Vision Zero approach to responsibility is new, and in a sense revolutionary, in traffic safety. However, it is certainly not without forerunners in other areas of safety. In a sense, it can be seen as the implementation in traffic safety of a general outlook that has long been taken for self-evident in workplace safety. In stark contrast to the traditional focus on individual fault and culpability in traffic safety, workplace safety has a strong and well-established focus on technological and organizational causal factors that can be eliminated or curtailed in order to reduce the prevalence of injuries. Since these factors are almost invariably in the employer’s control, it follows from this approach that the employer, rather than the employees, has the primary responsibility for safety on the workplace.

An interesting comparison can be made between the approaches to two types of traffic accidents, namely, road traffic accidents and accidents with forklift trucks on workplaces. As we have just noted, the road traffic literature still standardly looks for “the cause” of accidents and categorizes most accidents as caused by road users. In contrast, at least since the 1970s, the literature on forklift truck safety has refrained from looking for culprits and instead investigated the various types of forklift accidents with the purpose of “[p]rescribing the remedy (design improvement) to minimize the hazard and lower the risk” (MacCollum 1978, p. 145; cf. Stout-Wiegand 1987; Larsson and Rechnitzer 1994). One of the effects of Vision Zero is that the view of causality and responsibility that has since long been applied to forklifts, as well as to other dangerous machines on workplaces, is now increasingly applied to motor vehicles on public roads.

Self-Driving Cars

Self-driving cars have been discussed and to some extent developed at least since the 1950s, but it is only in the twenty-first century that they have become a realistic possibility. One of the several ethical issues that their potential introduction gives rise to is that of responsibility, in particular blame responsibility. Self-driving cars are predicted to be involved in fewer accidents than conventional cars, but there will still be accidents. Who should be held responsible for such accidents?

There are four reasonably plausible answers to this question. Blame responsibility for accidents can be assigned to:

- The car itself, or more precisely to the *artificial intelligence* built into it
- The *users* who travel in the cars
- *No one* at all, just like no one is held responsible for the occurrence of natural disasters
- *Other persons* than the users

Let us consider each of these options in turn. Concerning the first, it is important to distinguish between the question whether the artificial intelligence in self-driving cars can be held responsible for accidents and the much more general question whether any artificial intelligence can at all be held responsible in the same way as we hold human beings responsible for their doings (Nyholm 2018a, pp. 1209–1210). The answer to the latter question seems to depend crucially on what types of artificial intelligence humans will encounter in the future. We can think of hypothetical future intelligences that will exhibit beliefs and desires and communicate with humans about moral issues in much the same way that we humans communicate with each other. It is fairly plausible that we, or future humans, would be disposed to assign blame (and task) responsibilities to such artificial agents, if and when we encounter them. However, this is not the type of artificial intelligence that will be installed in self-driving cars. Instead, these vehicles will be provided with software that is constructed to execute the orders given by humans and to do so in accordance

with guidelines and restrictions devised by their human designers. Therefore, it seems highly unlikely that we will treat them as agents that can be culpable or held responsible (Brey 2013; Purves et al. 2015; Coeckelbergh 2016; Nyholm 2018b).

Our second option is to hold the users of self-driving cars blame responsible for whatever damage the car is deemed to cause. Concerning this option, it is important to distinguish between semi-automated and fully automated vehicles. A semi-automated vehicle still has a “driver,” who is passively but constantly following the driving and prepared to intervene immediately whenever necessary. With this arrangement, it does not appear unreasonable to assign blame responsibility to a (standby) driver who did not take over and solve a situation that the system could not solve. A fully automated vehicle does not require a standby driver. Such a car can navigate on the roads without any human driver or passenger or when everyone onboard is asleep. It is difficult to see how blame responsibility for an accident could be assigned to the occupants of a vehicle under such circumstances.

The third option is to refrain from assigning blame responsibility for an accident to anyone at all. This is how we often react to natural events. We do not assign blame responsibility to anyone for the occurrence of hurricanes or tsunamis (although we often assign blame to people who have failed to prepare properly for such events). However, this is not how we react to machines or other technological devices that are causal factors in an accident. As noted above, we have a strong tendency to focus on causal factors that involve human actions, whenever there are any such factors. For example, automated train systems have been introduced in many parts of the world, mostly in metro networks and airport transit systems (Wang et al. 2016). Automated trains are subject to extensive safety management. Accidents are certainly not treated in the same way as unevadable natural disasters. Instead, they are treated in the same way as accidents in other, less automated technological systems, namely, as avoidable failings for which human beings are responsible (Seng et al. 2009). There is no reason to believe that accidents in automated systems on roads will be treated differently.

This leaves us with the fourth option, namely, to assign blame responsibility for accidents to some other persons than the users of the automatic vehicles. There are strong reasons to assume that this is what is going to happen. The obvious candidates for undertaking responsibility are the system designers and system owners, i.e., those who are responsible for the construction of the vehicles and the construction, maintenance, and management of the roads and the communication systems that these vehicles will operate with. This is how we assign responsibilities for other automated systems, such as the automated trains just mentioned. Importantly, this is also how the first serious accidents involving self-driving cars have been dealt with. In media and in public discussions, the responsibility of the car manufacturers has been taken for granted. There are also clear signs that the automobile industry is planning to assume that responsibility (Atiyeh 2015; Nyholm 2018b).

In conclusion, we have strong reasons to expect that the blame responsibility for accidents implicating self-driving cars will be assigned to designers and owners of the automated traffic system. But that is only part of the answer to our question. Our

future traffic systems will have many designers and owners. If two self-driving cars of different brands collide, then responsibilities may have to be distributed among two automobile companies, the organization responsible for maintenance of the road, the organization(s) running the electronic communication system(s) that guided and coordinated the two vehicles, and various subcontractors of these companies and organizations. The automobile industry has a history of protracted blame games (Noggle and Palmer 2005), and neither legal battles nor public relations campaigns over these responsibilities should come as a surprise. Instead of the philosophically fascinating, but probably unrealistic, issue whether we can assign responsibility to the self-driving cars themselves, we may have to deal with more mundane conflicts between companies trying to avoid financial and reputational losses.

Institutional and Professional Responsibility

Much of the discussion above boils down to the unavoidable conclusion that in order to improve traffic safety, it is not sufficient to remind road users of their responsibilities. First and foremost, we have to assign important task responsibilities to those who have the resources and the power to bring about such improvements, namely, the system designers. This will not always be easy. Road traffic has no single responsible authority corresponding to the employer on a workplace. Instead, it has a large and rather heterogeneous collection of system designers.

Who then are the system designers? In the Swedish VZ policy, the concept embraces all actors—public and private—who, in their professional capacity, influence the design and function of the road system. . . . Three groups of designers were singled out as particularly important: road administrators (state, municipalities, and private), the automotive industry, and actors procuring or providing transport services (taxi, bus, and freight). Other identified system designers are actors responsible for various support systems, such as the police (monitoring and enforcement), driving schools (education), and emergency services, health care, and rehabilitation professionals. (Hysing 2021)

As yet, the responsibility of system designers is largely informal. Those working in the public sector are of course required to implement the government's policy, but the involvement of the automotive industry and other private sector companies and organizations is voluntary. Furthermore, there is no liability associated with these responsibilities. This can be compared to the employer's responsibility for workplace safety, which is in most jurisdictions fairly far-reaching and subject to legal sanctions. As discussed in Abebe et al. (2022), the lack of legal liability for system designers has been the object of some criticism, but it is not clear whether the introduction of such liability would lead to improvements in safety.

Importantly, the system designers' responsibility is a matter of both institutional and professional responsibility. The institutional responsibility is carried out by government agencies and private companies. The professional responsibility is carried out by the traffic safety experts who work in these institutions.

The notion of specific professional responsibilities goes back at least to the Greek physician Hippocrates (c. 460–c. 370 BCE), whose oath for physicians made it clear that a physician, when acting as a professional, has special duties and responsibilities that differ from those of citizens in general. In the Hippocratic tradition, the physician had to act for the benefit of the ill, keep silent about what he learnt about patients and their families, and treat all patients alike irrespective of whether they were men or women, rich or poor, free or slaves. If needed he should offer his service for free to the poor (Jouanna 1999, pp. 112–126; Askitopoulou and Vgontzas 2018). These are still parts of the ethics of the medical profession. However, there is also an important difference: Whereas the Hippocratic physician acted alone, physicians in the modern world work together with others. Instead of the single physician visiting patients in their homes, healthcare is now mainly performed in teams consisting of physicians and other healthcare personnel with different specializations (Heubel 2015).

What makes medical ethics *professional* is that it puts focus on certain values for which members of the profession have a special responsibility. For instance, according to the Hippocratic oath, the physician should always be of service to the ill. Therefore, he could not undertake to kill or hurt a person with a poison or suggest to others how to do so (Jouanna 1999, pp. 128–131). Interestingly, this was recognized by Plato, who considered it a more serious crime for a physician than for a layperson to poison a person. (Plato, *Laws*, Chap. 11, p. 933d.) Still today, all major organizations in the medical profession disallow their members to contribute in any way to capital punishment (Anon 2005; Litton 2013). In this context, the American Medical Association has made it very clear that professional ethics is distinct from personal moral judgments:

An individual's opinion on capital punishment is the personal moral decision of the individual. A physician, as a member of a profession dedicated to preserving life when there is hope of doing so, should not be a participant in an execution. (American Medical Association 2019)

Much later than physicians, other professions have developed professional identities, responsibilities, and ethical principles of their own. Lawyers, accountants, and engineers are among the most prominent examples. For instance, since the late nineteenth century, the engineering profession has developed ethical codes and delineated specific responsibilities that follow with the profession of an engineer. The value that has most often been associated with engineering professionalism is that of safety. Just as the ethical codes of physicians prevent them from undertaking to poison or otherwise hurt a person (even if it is legal to do so), the ethical codes of engineers prevent them from accepting assignments to make unsafe or dangerous constructions. Importantly, the ethical requirement not to compromise on safety is considered to override contractual obligations towards employers and customers.

Road safety has not yet been established as a profession like those of medicine and engineering, but the professionals whose work determines the risks we all run as road users can easily be identified. Although the overarching responsibilities for

traffic safety has to be assigned to the organizations that make and maintain roads and vehicles, the practical day-to-day implementation of these responsibilities will require the direct engagement of those who actually do the work. It is difficult to see how patient safety could be achieved in a hospital solely through directives from the top, without authorizing competent professionals to independently promote safety in their daily work. Road safety may not be very different from patient safety in this respect.

Cross-References

- ▶ [Arguments Against Vision Zero: A Literature Review](#)
- ▶ [ISO 39001 Road Traffic Safety Management System, Performance Recording, and Reporting](#)
- ▶ [Saving Lives Beyond 2020: The Next Steps](#)
- ▶ [Vision Zero in Sweden: Streaming Through Problems, Politics, and Policies](#)
- ▶ [Vision Zero: How It All Started](#)

Acknowledgments I would like to thank Claes Tingvall and Henok Girma Abebe for valuable comments on an earlier version of this chapter.

References

- Abebe, H. G., Edvardsson Björnberg, K., & Hansson, S. O. (2022). Arguments against Vision Zero: A literature review. This volume.
- Algora-Buenafé, A. F., Suasnavas-Bermúdez, P. R., Merino-Salazar, P., & Gómez-García, A. R. (2017). Epidemiological study of fatal road traffic accidents in Ecuador. *The Australasian Medical Journal*, 10(3), 238–245.
- American Medical Association. (2019). Code of Medical Ethics opinion 9.7.3. <https://www.ama-assn.org/delivering-care/ethics/capital-punishment>. Downloaded 28 July 2020.
- Anon. (2005). Medical collusion in the death penalty: An American atrocity [editorial]. *Lancet*, 365, 1361–1361.
- Anon. (2013). Transport safety award open for entries. <https://cvshow.com/transport-safety-award-open-for-entries/>. Downloaded 23 Feb 2020.
- Askitopoulou, H., & Vgontzas, A. N. (2018). The relevance of the Hippocratic Oath to the ethical and moral values of contemporary medicine. Part II: Interpretation of the Hippocratic Oath – Today’s perspective. *European Spine Journal*, 27(7), 1491–1500.
- Atiyeh, C. (2015). Volvo will take responsibility if its self-driving cars crash. *Car and Driver*, October 8. <https://www.caranddriver.com/news/a15352720/volvo-will-take-responsibility-if-its-self-driving-cars-crash/>
- Belin, M.-Å., Tillgren, P., & Vedung, E. (2012). Vision Zero – A road safety policy innovation. *International Journal of Injury Control and Safety Promotion*, 19(2), 171–179.
- Berger, R. (1998). Understanding science: Why causes are not enough. *Philosophy of Science*, 65, 306–332.
- Bernstein, S. (2017). Causal proportions and moral responsibility. In D. Shoemaker (Ed.), *Oxford studies in agency and responsibility* (Vol. 4). Oxford, UK: Oxford University Press.

- Boyd, M. (2015). A method for prioritizing interventions following root cause analysis (RCA): Lessons from philosophy. *Journal of Evaluation in Clinical Practice*, 21(3), 461–469.
- Brey, P. (2013). From moral agents to moral factors: The structural ethics approach. In P. Kroes & P.-P. Verbeer (Eds.), *The moral status of artifacts* (pp. 125–142). Dordrecht: Springer.
- Cane, P. (2002). *Responsibility in law and morality*. Oxford, UK: Hart Publishing.
- Coeckelbergh, M. (2016). Responsibility and the moral phenomenology of using self-driving cars. *Applied Artificial Intelligence*, 30(8), 748–757.
- Dent, E. B. (2003). The interaction model: An alternative to the direct cause and effect construct for mutually causal organizational phenomena. *Foundations of Science*, 8, 295–314.
- Descartes, R. ([1632] 1987). *Traité de l'Homme*. In R. Descartes (Ed.), *Oeuvres et lettres. Textes présentés par André Bridoux*. Paris: Gallimard.
- Duff, R. A. (1998). Responsibility. In *Routledge encyclopedia of philosophy*. London: Taylor and Francis. <https://www.rep.routledge.com/articles/thematic/responsibility/v-1>
- Dworkin, G. (1981). Voluntary health risks and public policy: Taking risks, assessing responsibility. *The Hastings Center Report*, 11, 26–31.
- Goodin, R. E. (1987). Apportioning responsibilities. *Law and Philosophy*, 6, 167–185.
- Government Bill 1996/97:137. (1996). *Nollvisionen och det Trafiksäkra Samhället [Vision Zero and the traffic safe society]*. Stockholm: Swedish Government.
- Grill, K., & Fahlquist, J. N. (2012). Responsibility, paternalism and alcohol interlocks. *Public Health Ethics*, 5(2), 116–127.
- Hansson, S. O. (2007). Hypothetical retrospection. *Ethical Theory and Moral Practice*, 10, 145–157.
- Hansson, S. O. (2010). The harmful influence of decision theory on ethics. *Ethical Theory and Moral Practice*, 13, 585–593.
- Hansson, S. O. (2018). Politique du risque et l'intégrité de la science. In A. de Guay (Ed.), *Risque et Expertise. Les Conférences Pierre Duhem* (pp. 57–112). Besançon: Presses Universitaires de Franche-Comté.
- Hansson, S. O. (2022). Who is responsible if the car itself is driving?. In Diane P. Michelfelder, *Test-driving the future*. Rowman & Littlefield, in press.
- Hansson, S. O., & Peterson, M. (2001). Rights, risks, and residual obligations. *Risk Decision and Policy*, 6, 157–166.
- Harantová, V., Kubiková, S., & Rumanovský, L. (2019). Traffic accident occurrence, its prediction and causes. In J. Mikulski (Ed.), *Development of transport by telematics. TST 2019* (Communications in computer and information science. Vol. 1049, pp. 123–136). Cham: Springer.
- Hart, H. L. A. (2008). *Punishment and responsibility. Essays in the philosophy of law, 2nd edn, with an introduction by John Gardner*. Oxford, UK: Oxford University Press.
- Hausman, D. M., & Woodward, J. (1999). Independence, invariance and the causal Markov condition. *British Journal for the Philosophy of Science*, 50, 521–583.
- Heubel, F. (2015). The “soul of professionalism” in the Hippocratic Oath and today. *Medicine, Health Care and Philosophy*, 18(2), 185–194.
- Honoré, A., & Gardner, J. (2010). Causation in the law. In E. N. Zalta (Ed.), *Stanford encyclopedia of philosophy*. Downloaded 23 Nov 2019 from <https://plato.stanford.edu/archives/fall2019/entries/causation-law/>
- Howarth, I. (1985). Interactions between drivers and pedestrians: Some new approaches to pedestrian safety. In L. Evans & R. C. Schwing (Eds.), *Human behavior and traffic safety* (pp. 171–178). New York: Plenum Press.
- Howarth, C. I., & Repetto-Wright, R. (1978). The measurement of risk and the attribution of responsibility for child pedestrian accidents. *Safety Education*, 144, 10–13.
- Huber, S. J., & Wynia, M. K. (2004). When pestilence prevails ... physician responsibilities in epidemics. *American Journal of Bioethics*, 4(1), 5–11.
- Husak, D. (2004). Vehicles and crashes: Why is this moral issue overlooked? *Social Theory and Practice*, 30(3), 351–370.

- Hydén, C. (2019). Speed in a high-speed society. *International Journal of Injury Control and Safety Promotion*. <https://doi.org/10.1080/17457300.2019.1680566>.
- Hysing, E. (2021). Responsibilization: The case of road safety governance. *Regulation and Governance*, 15, 356–369.
- Jouanna, J. (1999). *Hippocrates*. Baltimore: Johns Hopkins University Press.
- Kuhn, T. S. (1971). La notion de causalité dans le développement de la physique. In M. Bunge (Ed.), *Les Théories de la Causalité* (pp. 4–15). Paris: Presses Universitaires de France.
- Larsson, T. J., & Rechnitzer, G. (1994). Forklift trucks – Analysis of severe and fatal occupational injuries, critical incidents and priorities for prevention. *Safety Science*, 17(4), 275–289.
- Lewis, D. (1973). Causation. *Journal of Philosophy*, 70(17), 556–567.
- Litton, P. J. (2013). Physician participation in executions, the morality of capital punishment, and the practical implications of their relationship. *The Journal of Law, Medicine & Ethics*, 41, 333–352.
- MacCollum, D. V. (1978). How safe the lift? *Proceedings of the Human Factors Society Annual Meeting*, 22(1), 145–149.
- Michaud, P.-A., Blum, R. W., Benaroyo, L., Zermatten, J., & Baltag, V. (2015). Assessing an adolescent's capacity for autonomous decision-making in clinical care. *Journal of Adolescent Health*, 57(4), 361–366.
- Mill, J. S. ([1843] 1996). A system of logic. In: Robson, J. M. (Ed.), *Collected works of John Stuart Mill* (Vol. 7). Toronto: University of Toronto Press.
- Moore, M. S. (1999). Causation and responsibility. *Social Philosophy and Policy*, 16(2), 1–51.
- Moore, M. S. (2009). *Causation and responsibility: An essay in law, morals, and metaphysics*. New York: Oxford University Press.
- Mumford, S. (2013). Causes for laws. *Jurisprudence*, 4(1), 109–114.
- Nilsson, H. (2020a). Just nu: Volvos strategi hotar exporten [Just now: Volvo's strategy threatens the export]. Dagens PS, April 22. <https://www.dagensps.se/nyheter/volvos-strategi-hotar-exporten/>. Downloaded 28 July 2020.
- Nilsson, H. (2020b). Läsarna rasar: “Volvo gräver sin egen grav” [The readers are furious: Volvo is digging its own grave]. Dagens PS, April 24. <https://www.dagensps.se/motor/bilteknik/lasarna-rasar-volvo-graver-sin-egen-grav/>. Downloaded 28 July 2020.
- Noggle, R., & Palmer, D. E. (2005). Radials, rollovers and responsibility: An examination of the Ford-Firestone case. *Journal of Business Ethics*, 56(2), 185–204.
- Nyholm, S. (2018a). Attributing agency to automated systems: Reflections on human–robot collaborations and responsibility-loci. *Science and Engineering Ethics*, 24(4), 1201–1219.
- Nyholm, S. (2018b). The ethics of crashes with self-driving cars: A roadmap, II. *Philosophy Compass*, 13(7), e12506.
- Ori, M. (2014). The morality of motorcycling. *Philosophical Papers*, 43(3), 345–363.
- Parker, M. (2001). The ethics of evidence-based patient choice. *Health Expectations*, 4(2), 87–91.
- Parry, G. W. (1991). Common cause failure analysis: A critique and some suggestions. *Reliability Engineering & System Safety*, 34(3), 309–326.
- Purves, D., Jenkins, R., & Strawser, B. J. (2015). Autonomous machines, moral judgment, and acting for the right reasons. *Ethical Theory and Moral Practice*, 18(4), 851–872.
- Reiff, M. R. (2015). No such thing as accident: Rethinking the relation between causal and moral responsibility. *Canadian Journal of Law and Jurisprudence*, 28(2), 371–397.
- Rizzi, D. A., & Pedersen, S. A. (1992). Causality in medicine: Towards a theory and terminology. *Theoretical Medicine*, 13, 233–254.
- Rooney, J. J., & Vanden Heuvel, L. N. (2004). Root cause analysis for beginners. *Quality Progress*, 37(7), 45–56.
- Rosencrantz, H., Edvardsson, K., & Hansson, S. O. (2007). Vision Zero – Is it irrational? *Transportation Research Part A: Policy and Practice*, 41, 559–567.
- Russell, B. (1913). On the notion of a cause. *Proceedings of the Aristotelian Society*, 13, 1–26. Reprinted in Bertrand Russell, B. (1994). *Mysticism and logic* (pp. 173–199). London: Routledge.

- Santosa, S. P., Mahyuddin, A. I., & Sunoto, F. G. (2017). Anatomy of injury severity and fatality in Indonesian traffic accidents. *Journal of Engineering and Technological Sciences*, 49(3), 412–422.
- Seng, Y. K., Wai, N. H., Chan, S., & Weng, L. K. (2009). Automated metro – Ensuring safety and reliability with minimum human intervention. *INCOSE International Symposium*, 19(1), 1–10.
- Shaver, K. G. (1985). *The attribution of blame. Causality, responsibility, and blameworthiness*. New York: Springer.
- Stirrat, G. M., & Gill, R. (2005). Autonomy in medical ethics after O'Neill. *Journal of Medical Ethics*, 31(3), 127–130.
- Stout-Wiegand, N. (1987). Characteristics of work-related injuries involving forklift trucks. *Journal of Safety Research*, 18(4), 179–190.
- Tarnovanu, H. (2015). Causation and responsibility: Four aspects of their relation (Ph.D. thesis). University of St. Andrews.
- van de Poel, I. (2011). The relation between forward-looking and backward-looking responsibility. In N. A. Vincent, I. van de Poel, & J. van den Hoven (Eds.), *Moral responsibility* (pp. 37–52). Dordrecht: Springer.
- Verweij, M., & Dawson, A. (2019). Sharing responsibility: Responsibility for health is not a zero-sum game. *Public Health Ethics*, 12(2), 99–102.
- Walster, E. (1966). Assignment of responsibility for an accident. *Journal of Personality and Social Psychology*, 3(1), 73–79.
- Wang, Y., Zhang, M., Ma, J., & Zhou, X. (2016). Survey on driverless train operation for urban rail transit systems. *Urban Rail Transit*, 2(3–4), 106–113.
- Wetmore, J. M. (2004). Redefining risks and redistributing responsibilities: Building networks to increase automobile safety. *Science, Technology & Human Values*, 29(3), 377–405.
- Willemsen, P., & Kirfel, L. (2019). Recent empirical work on the relationship between causal judgements and norms. *Philosophy Compass*, 14(1), e12562.
- Wittgenstein, L. (1976). Cause and effect: Intuitive awareness. *Philosophia*, 6(3), 409–425.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.



Liberty, Paternalism, and Road Safety

6

Sven Ove Hansson

Contents

Introduction	206
Paternalism and Liberty	207
Defining the Term	207
How Bad Is Benevolence?	210
What Options Can We Forbear?	213
Sanitation	213
Workplace Safety	214
Seat Belts	215
Helmets	217
Drunk Driving	219
Speeding	221
Conclusions	223
The Significance of Human Connectedness	225
Combined Causes and Extended Anti-paternalism	226
Herd Effects: How We All Influence Each Other	230
Driver Assistance and Self-driving Cars	233
In Conclusion: Vision Zero	235
Cross-References	236
References	236

Abstract

Traffic safety measures such as seat belts, helmets, and speed limits have often been opposed by people claiming that these measures infringe on their liberty. Safety measures are often described as paternalistic, i.e., as protecting people against their own will. This chapter provides a historical account of such criticism of safety measures, beginning with nineteenth-century opposition to sanitation

S. O. Hansson (✉)

Division of Philosophy, Royal Institute of Technology (KTH), Stockholm, Sweden

Department of Learning, Informatics, Management and Ethics, Karolinska Institutet, Stockholm, Sweden

e-mail: soh@kth.se

© The Author(s) 2023

K. Edvardsson Björnberg et al. (eds.), *The Vision Zero Handbook*,
https://doi.org/10.1007/978-3-030-76505-7_6

205

measures, which were claimed to threaten the freedom to drink dirty water. The historical analysis has a surprising conclusion: Opposition to safety measures does not seem to have much to do with paternalism. Some measures that would typically be described as paternalistic, such as seat belts in commercial aviation and hard hats on construction sites, have met with no significant opposition. In contrast, some of the most vehemently opposed measures, such as speed limits and the prohibition of drunk driving, cannot with any vestige of credibility be described as paternalistic. This is followed by an analysis showing that due to our tendency to follow examples set by others (herd effects), purely self-affecting behavior is much less common than what has usually been assumed. Most of the opposition to safety measures in road traffic seem to result from some individuals' desires to engage in activities that endanger other people's lives. The social need to restrain the satisfaction of such desires is obvious.

Keywords

Acceptance of safety measures · Herd effects · Liberty · Paternalism · Traffic safety · Vision Zero

Introduction

In April 1985, the Senate of the State of Washington discussed a proposal to make the use of seat belts mandatory. Arguing against the proposal, Senator Kent Pullen (R) conceded that seat belts would save lives, but in his view, that was not reason enough to make them mandatory. "There is something more important than life itself," he said, "and that's freedom" (Leichter 1991, p. 12). This is a type of argument that traffic safety has met with throughout its history, and it is still in very active use. Currently, a Missouri law-maker, Eric Burlison (R), is working hard to repeal the motorcycle helmet law in his state. Like Pullen, he recognizes that his policies will lead to more deaths, but that does not deter him. "At the end of the day," he says, "it's about individual responsibility and individual freedom. I want my neighbor to stay safe and healthy, but it's not my business to force those decisions upon my neighbor" (Huguelet 2020).

We hear similar messages in many other contexts. The American Enterprise Institute campaigns against the "spirit of anti-smoking paternalism" that has given us smoke-free bars and restaurants, and they deplore attempts by legislators to remove unhealthy components such as trans-fats from food, calling all this "a remarkable confiscation of freedom" (White 2006). The owner of a British pub who was criticized by the authorities for deficient hygiene complained about "our modern nanny state's requirements for sterile and salubrious certification" (Woolfson 2016). Anti-vaccination campaigners call themselves "freedom keepers," and claim that parent's rights not to have their children vaccinated is a matter of "civil rights" (Mays 2019). And in 2017, an Australian parliamentarian declared the prevention of such measures to be his main aim in politics.

When our most basic rights and freedoms are being chipped away at on a daily basis through nanny-state regulations and big-government paternalism, with smoking indoors banned, irrespective of what the owner of the property thinks; with bicycle helmets mandatory, despite the rest of the world agreeing that they are really not worth the effort; and with e-cigarettes, a potentially life-saving alternative pathway to quitting smoking, practically banned, I will be there, making the case for personal choice and personal responsibility. (Stonehouse 2017)

What is all this about? Of course we want freedom, but we also want traffic safety, healthy food, and protection against deadly diseases. How severe is the conflict between safety and liberty? Can we have safety on the roads in a free society, or do we have to choose between freedom and safety?

Paternalism and Liberty

At the center of this discussion is the concept of *paternalism*, protecting someone against her or his will. The critics of government interventions – or at least the more thoughtful among them – accept government interventions to prevent people from harming others, but they reject interventions intended to prevent us from harming ourselves. The latter type of regulation, they say, treats adults as if they were children. That is why it is called paternalistic.

Defining the Term

The word “paternalism” is derived from “paternal,” which means fatherly. In the late nineteenth century, it was often used in a positive sense, in particular about employers who cared for their employees. Today it is almost universally used in a negative sense. In the scholarly discussion, Gerald Dworkin’s definition of paternalism is widely accepted. In its latest variant it says:

Paternalism is the interference of a state or an individual with another person, against their will, and defended or motivated by a claim that the person interfered with will be better off or protected from harm. (Dworkin 2020)

Three major components of this definition should be noticed: (1) a person is interfered with, (2) the person in question does not desire that intervention, and (3) the interference has the purpose to benefit this person (Wilson 2015).

A small modification and some clarification of the definition will be useful. To begin with the modification, Dworkin only mentions that “a state or an individual” can interfere with a person, but certainly so can also a company or some other type of organization. For instance, a company that performs construction work may choose to fence in the worksite for no other reason than a desire to protect the public from danger. Similarly, a manufacturer may choose to voluntarily withdraw a product because of problems with its quality or safety, although some consumers still want

the product. A company that takes such actions behaves paternalistically in the same sense as a government that makes the same type of decisions for the same type of reasons. In these and other cases, an action affecting a group of persons can be paternalistic towards some but not all of its members, depending on whether or not they oppose the action (Grill 2018).

Dworkin's term "interfere with" is in need of clarification, and indeed he provides such a clarification later on in the text, when he specifies that the paternalist "interferes with the liberty or autonomy" of the targeted person (Dworkin 2020). This is an important clarification, since otherwise the list of paternalistic actions would grow to enormous proportions. In private life and in work-life we often find ourselves doing something for the benefit of a person who does not appreciate it. And unavoidably, political decisions intended to benefit large groups of people are not always appreciated by all of them. Some writers have used the term "paternalism" for all sorts of political decisions that are intended to benefit the population, including government funding of healthcare, education, public broadcasting, museums and theaters, tax deductions for pension savings, and subsidies of leisure activities and healthy food (Le Grand and New 2015, pp. 12, 52–54, 65–66, and 153; Conly 2017, p. 208). With such a wide definition, the building and maintenance of roads is also paternalistic, and so most certainly is law enforcement, which – just like publicly funded art museums – is intended to benefit all, but unfortunately is not appreciated by all. If we call all of this paternalism, we risk losing sight of the much more limited range of cases that have been at the focus of debates on paternalism, such as seat belts, motorcycle helmets, and vaccination. These cases all have in common that they satisfy Dworkin's criterion of interfering with the liberty or autonomy of the targeted person (Cf. also Dworkin 1972, p. 67). Another way to express this is that they limit the person's choices by removing an option, making an option less accessible, or reducing her ability to choose.

In definitions and more principled discussions of paternalism it is assumed that a person has a right to harm or risk her own health and well-being, but it is not assumed that she also has a right to harm her family members or expose them to risk. However, in political discussions, and even in some academic texts, the latter assumption is made as well. Even Giubilini and Savulescu (2019, p. 241), who commendably advocate universal vaccination programs, describe this as a matter of "how much weight we want to give to paternalism compared to individual freedom, and particularly the freedom to decide what kind of risks to take on oneself or on one's children." Shiffrin (2000, p. 217) takes this even further, suggesting that a park ranger behaves paternalistically if he prevents a person from making a dangerous mountain climb, not at all out of concern for the person himself or herself but for the person's spouse who would be left grief-stricken in the case of a fatal accident. As noted by Dworkin (2015, p. 26), according to how paternalism is usually defined, this case rather exemplifies "precisely the contrast class to paternalism."

Non-paternalistic measures to promote safety and public health have a long tradition. For instance, free vaccinations were an important part of infection control

throughout the nineteenth century, and they still are (Rigau-Pérez 1989, p. 400; Daker et al. 1893, p. 217; Anon 1894). In order to encourage traffic behavior that reduces the risk of crashes, stop lines have been painted on roads since 1907, and centerlines and pedestrian crossings since 1911 (Petroski 2016). The strategy behind these and other similar measures for health and safety was summarized in 1976 by the American public health scholar Nancy Milio:

The point is that most human beings, professional or nonprofessional, provider or consumer, make the easiest choices available to them most of the time, and not necessarily because of what they know is most healthful. Thus, if it is agreed that health-promoting life patterns are a good thing, then the focus for changing behavior should be on the problem of how to make health-generating choices more easy, and how to make health-damaging choices more difficult. . .

In order then for life-style patterns to alter among individuals in numbers sufficient to affect the incidence of major diseases, new, health-promoting options must be available, and more readily so than health-damaging options, i.e., in such a way as to be less costly in dollar and other costs. (Milio 1976, pp. 435 and 437)

In 2008, the term “nudge” was introduced in a book promoting this approach (Thaler and Sunstein 2008). Neither this nor most other texts on “nudge” acknowledges the long tradition of similar ideas in public health. This was pointed out by Signild Vallgård (2012), who also noted that “while renaming an activity has the benefit of creating enthusiasm and new energy and providing the encouraging feeling of being part of something new, it also has a drawback: it could lead to a neglect of insights generated by the pursuance of similar policies in the past.” In their book, Thaler and Sunstein used the term “libertarian paternalism” about the “nudge” approach. As already mentioned, this is a usage that, if applied consistently, would lead us to classify a large part of the benevolent activities that take place in society as “paternalistic.” It should be recognized, though, that some of the actions that are “libertarian paternalistic” according to Thaler and Sunstein, but not paternalistic according to traditional definitions, are problematic in other ways, for instance, by being manipulative (Mols et al. 2015; Dworkin 2020). This was also noted by Dworkin, who said:

Their definition of Paternalism is very weak in the sense that it allows many more acts to count as paternalistic than would be under almost all traditional definitions of paternalism. (Dworkin 2020)

In order to avoid thinning out the concept, I will use the term “paternalism” only for interventions that satisfy Dworkin’s definition, as explicated above. This is the terminology that appears to be most in line with the tradition, at least in philosophy. Of course, a different terminology could have been chosen. We could choose to use the term “paternalistic” about all measures that have the purpose to benefit a person against that person’s will, or without her consent. Then, however, we would have to introduce a new term for what is called “paternalism” here.

How Bad Is Benevolence?

Dworkin's definition of paternalism has a peculiar feature, which it apparently shares with all other definitions of the term. This is his third criterion, which says that a paternalistic action is "defended or motivated by a claim that the person interfered with will be better off or protected from harm." We can call this the criterion of benevolence.

Something seems to be wrong here. That an action affecting another person is benevolent towards that person is surely a morally good feature of that action. But paternalism is presumably bad. So how can a good feature of an action be a necessary requirement for the action to belong to a category of bad actions?

The mystery deepens if we consider how benevolence is usually assessed from a moral point of view. Consider the following example (which happens to be based on an actual occurrence):

After spreading thawing salt (road salt) on the asphalt outside his own house, Ahmad walked over to his elderly neighbor's house and strewed some salt on her entrance stairs. When she noticed this a couple of hours later, she was much annoyed since she believes that thawing salt will cause the stairs to crack.

Case A: Ahmad did this because he was worried that she might otherwise slip and break her leg.

Case B: Ahmad did this because his wife had told him that another neighbor wanted him to sprinkle salt over her stairs, but he mistakenly went to the wrong neighbor's house.

Case C: Ahmad did this because he wanted to wreck her stairs.

In case A, Ahmad's action is benevolent. In case C it is malevolent. In case B it is neither benevolent nor malevolent; we can call it neutral in that case.

It can, I believe, be safely assumed that most of us would agree that Ahmad's action is morally better in case B than in case C, and even better in case A. Wanting to help your neighbor seems to be a better excuse than confusing her with someone else. Case A is the case in which Ahmad's action was benevolent. Apparently, this example confirms the supposition that benevolence contributes to making actions good, rather than to making them bad.

But perhaps this example is untypical since it concerns actions by individuals, rather than actions by governments or organizations? Let us consider such examples, and begin with an example of a business organization.

Liliana owns the only hardware store in town. She has stopped selling a particular brand of electric jigsaws, although several customers would still like to buy one.

Case A: Liliana did this because the protection of the user's hands on this particular saw is inferior to that of other brands, and she did not want her customers to be hurt.

Case B: Liliana did this because her supplier did not offer this particular saw any more, and she did not take the trouble to look for someone else who can supply it.

In case A, Liliana's purpose is clearly benevolent towards the customers whom she has deprived of the option to buy this particular saw. Furthermore, according to standard usage of the term, her decision in this case is paternalistic. In case B, in

contrast, her action is neither benevolent nor paternalistic. It is just an ordinary business decision. The decision in case A is clearly praiseworthy from a moral point of view, and we might even use it as an example of corporate social responsibility on a small scale. In case B, her decision does not seem to be particularly praiseworthy from a moral point of view. So is paternalism better than business as usual? Then it cannot be terribly bad, or can it?

Let us consider a similar example where the government is involved:

Liliana had to stop selling the power jigsaw. It would have been illegal to sell it since the government did not prolong its type approval.

Case A: The government decided not to prolong the type approval because the jigsaw did not have a satisfactory protection of the user's hands.

Case B: The reason why the type approval was not prolonged was that a government official failed to include this brand on a list of electric tools for routine prolongation of the type approval. Due to procedural rules introduced by a previous government, it will take a full year to correct the mistake.

In case A, the motive of the decision was benevolent, and it can also be described as paternalistic. And again, the paternalistically motivated version of the decision would appear to be the morally best version.

Thus, it seems as if benevolence cannot make an action or a decision worse than what it would otherwise have been. Then how can benevolence be a defining characteristic of an undesirable feature of an action, namely, that it is paternalistic? Can the notion of paternalism at all be saved from this conundrum?

Yes, there is in fact a fairly satisfactory solution to this, but it requires a small change in the definition: The definition should not refer to the presence of paternalistic justifications of the action, but rather to the absence of sufficient non-paternalistic justifications. This is also what Dworkin indicates. When he elaborates his benevolence criterion more in detail, the crucial phrase is:

X does so only because X believes Z will improve the welfare of Y (Dworkin 2020; emphasis added)

(*X* is the agent who behaves paternalistically, *Z* the paternalistic action, and *Y* the person who is the target of the paternalistic action.)

In other words, a sensible anti-paternalist cannot maintain that it is wrong to have paternalistic motives or justifications for one's actions. (The terminology is somewhat intricate. A motive or justification for an intervention affecting a person is commonly called "paternalistic" if it is benevolent towards that person, irrespectively of whether or not the intervention infringes on that person's liberty. This usage will be followed here.) The wrong will have to be identified as that of not having sufficient non-paternalistic motives or justifications. Let us consider an example that illustrates this:

For many years Stephen has earned his living by travelling around with a pendulum ride, which he sets up on various festivals and fairs. One day he is visited by a safety inspector

from a government agency. After discovering a serious fault in the machine, the inspector issues an injunction prohibiting Stephen from offering any more rides on it.

Case A: The inspector does this because of a serious risk that the attendant, that is Stephen himself, can be killed in an accident. Riders are not at risk.

Case B: The inspector does this because of a serious risk of an accident in which both the riders and the attendant can get killed.

Case C: The inspector does this because of a serious risk that the riders can be killed in an accident. The attendant is not at risk.

Stephen, we may assume, is a fervent anti-paternalist. In case A, he can legitimately claim that the inspector has acted against him on purely paternalistic grounds. Such a claim would not be tenable in the other two cases, since the risk to the riders is reason enough for the inspector's decision. But would it be reasonable for Stephen to claim that the inspector's argument for closing down his machine is weaker in case B than in case C, since a concern for Stephen's own well-being is involved in the former case but not in the latter?

INSPECTOR: I have decided to close down your machine because there is a large risk that riders can get killed if you continue to operate it.

STEPHEN: I fully respect your decision.

INSPECTOR: I appreciate that. I should also tell you that you would probably be killed yourself in such an accident, and of course that contributed to my decision.

STEPHEN: What are you saying? Are you a paternalist wanting to protect me? Then I cannot respect your decision any longer. Are you really sure that it is necessary to close down the machine immediately?

It would not be absurd to claim that a paternalist argument has no weight at all. (This would mean that the inspector's injunction is justified to the same degree in cases B and C.) However, as this example shows, it would be utterly absurd to maintain that a paternalist argument carries a negative weight. There may of course be anti-paternalists who maintain that it is blameworthy to harbor concerns for other people's well-being. I am not claiming that such a position is impossible or non-existent, only that is a morally absurd position that is not worth taking seriously. This is written during the Covid-19 epidemic in 2020. Among the many statements that have been made on what governments should and should not do in relation to this health crisis, I have as yet not heard a claim that governments have no business to be concerned with the population's health. (This would mean that the injunction is less well justified in case B than in case C.) This confirms the limit to sensible anti-paternalism that was proposed above: It has to be concerned with the absence of sufficient non-paternalistic motives or justifications, not with the presence of paternalistic ones. This also means that the legitimate concerns of anti-paternalism are concerns about infringements of liberty, not about other people's benevolence. Wilson (2015, p. 212) reached a similar conclusion. Several authors have pointed at the difficulties involved in judging an action by a government or an institution by its intention (e.g., Dworkin 1972, p. 65). Participants in a decision may differ in their intentions, and it is far from clear how their intentions can be aggregated (Preyer et al. 2014; O'Madagain 2014; Salice 2015). I will give the last word to the safety inspector:

INSPECTOR: Look here, Stephen. I have full respect for your strong views on personal liberty. I certainly cherish the idea of a free society myself. But there is an imminent risk to your customers, and that is reason enough to close down the ride. The freedom that I believe in is not a freedom to hurt others. And if you don't like that I care for your safety, then let me just tell you one thing: I am in my right to do so. How can you bother about your own liberty without accepting my liberty to care and worry about other people? Would you now please close down the ride?

What Options Can We Forbear?

We have identified the bad that anti-paternalists legitimately worry about as infringements in liberty. Liberty consists in being able to choose between different options. We can therefore express this as a concern that people are deprived of options that they could otherwise have chosen among. This is clearly an important concern. Everyone's right to make their own decisions and to choose their own way of life is a crucial part of what it means to live in a free society.

Living in a society with other people comes with a wealth of options that an eremite living in the wilderness does not have. (Just think of what you have been doing today, before reading this. How much of that could you have done without the combined effects of the actions of innumerable other humans, of present and past generations?) As our societies develop, some options are lost while others are added. Some options have been lost due to commercial decisions; for instance, you cannot buy a new car with a 20 hp engine any more. (That was the power of the Ford Model T motor.) Other options have been lost due to government decisions; you cannot spray DDT on your garden roses to get rid of insects. (This use of the pesticide was still recommended by the US Department of Agriculture in 1967, 5 years after Rachel Carson's book *Silent Spring* (1962) was published (Smith 1967, pp. 6 and 9–12).) Both these are examples of lost options that few would wish to regain. They illustrate that not all losses of options are regrettable. In the case of DDT, there was considerable resistance to the change (Krupke et al. 2007), but that reaction has since long abated. Let us have a look at some examples of safety-related infringements of liberty, what reactions they have encountered and how those reactions have developed.

Sanitation

Our first example is a classic issue in public health, namely, clean water and hygienic sewage systems. The British Public Health Act, adopted in 1848, authorized local authorities to take control over water supply, sewerage, and other facilities needed to ensure hygienic living conditions. The new law was strongly supported by those whose living conditions it would improve, as well as by philanthropists and radicals in higher social strata, but it also met with considerable resistance (Roberts 1958, 1979, pp. 200 and 258–259; Wiggins 1987; Porter 1999, pp. 119–120).

Speaking in the House of Commons, the Tory MP Charles Newdegate (1816–1887) described the law as “a departure from the free principles of the British Constitution.” His liberal colleague George Muntz (1794–1857) considered the whole issue of sanitation to be “grossly exaggerated,” and denounced the legislation as unnecessary since “[t]he people were clever enough to manage their own affairs” (House of Commons 1847, cc. 729 and 750). Another liberal MP, Edward Divett (1797–1864), warned against “the constant meddling and dictation” that he expected from the government authority overseeing the legislation. “The electors did not object, it was true,” he said, “to sanitary improvement; but they did not choose to be ordered how to set about it.” (House of Commons 1848, c. 725) The most prominent opponent of the law was the influential conservative MP David Urquhart (1805–1877), who opposed the bill because it placed “despotic powers” in the hands of the Government. Government had already too much power, he said, and he was determined to do whatever he could “not merely to prevent any increase to that power, but to reduce the amount of it which the Government at present enjoyed.” (House of Commons 1848, cc. 712 and 715)

The most vociferous critics could be found in the Tory papers. For instance, the London journal *Morning Herald* argued that “[a] little dirt and freedom may after all be more desirable than no dirt at all and slavery” (Roberts 1979, p. 200). However, the prevalence of such sentiments receded in the coming decades, and today we do not hear much talk about the importance of being free to drink contaminated water or walk on filthy streets.

Workplace Safety

The legal approach to workplace health and safety underwent large changes in the late nineteenth and early twentieth centuries. According to the old tradition, male workers were responsible for their own safety, and any attempt to lay responsibility on the employer was conceived as an infringement on the workmen’s freedom of contract. For instance, in 1880 a liberal government in Britain introduced new legislation that gave workers right to compensation from the employer for injuries incurred on the workplace. The Tories vehemently opposed this measure. In Parliament, the conservative member Thomas Knowles (1824–1883), who had business interests in coal mining, warned that the new legislation would be disastrous for mine owners, whose companies in his view had so many accidents that the law would in practice “confiscate their property.” He also told Parliament that due to the new liability legislation, he estimated that his own assets in coal mines had only half the value that they had had a week ago (Knowles 1880, cc. 1100–1101. Cf. Green 1995, p. 73.). There was also considerable resistance against the law within the liberal party (Green 1995, p. 75). The prominent left liberal philosopher Thomas Hill Green (1836–1882), who supported the legislation, summarized the most prominent argument of his opponents as follows:

“The workman,” it was argued, “should be left to take care of himself by the terms of his agreement with the employer. It is not for the state to step in and say, as by the new act it says, that when a workman is hurt in carrying out the instructions of the employer or his foreman, the employer, in the absence of a special agreement to the contrary, shall be liable for compensation. If the law thus takes to protecting men, whether tenant-farmers, or pitmen, or railway servants, who ought to be able to protect themselves, it tends to weaken their self-reliance, and thus, in unwisely seeking to do them good, it lowers them in the scale of moral beings.” (Green 1888 [1881], p. 365)

Crossley (1999, p. 292) quotes this passage, but removes Green’s four quotation marks as well as the phrase “it was argued”. He correctly attributes the text to Green, but omits the crucial information that Green described his opponents’ views, and incorrectly describes the text as written by “[o]ne opponent of the Act”. This, said Green, was an argument “which many of us, without being convinced by it, may have found it difficult to answer.” (Green, p. 365) But he had an answer:

When we speak of freedom as something to be so highly prized, we mean a positive power or capacity of doing or enjoying something worth doing or enjoying, and that, too, something that we do or enjoy in common with others. We mean by it a power which each man exercises through the help or security given him by his fellow-men, and which he in turn helps to secure for them. . . . If the ideal of true freedom is the maximum of power for all members of human society alike to make the best of themselves, we are right in refusing to ascribe the glory of freedom to a state in which the apparent elevation of the few is founded on the degradation of the many, and in ranking modern society, founded as it is on free industry, with all its confusion and ignorant licence and waste of effort, above the most splendid of ancient republics. (ibid., pp. 371–372)

A contract in which someone “bargains to work under conditions fatal to health” could not, in his view, be defended in the name of freedom since it was “an impediment to the general freedom,” preventing the workers from making the best of themselves (ibid., p. 373). At the core of his argument was his assertion that the freedom to take a dangerous job, without even a right to compensation in the case of an accident, was simply not a freedom worth having. This was also the general viewpoint among workers, and it was therefore Thomas Knowles and other opponents of this “paternalist” legislation, rather than its proponents, who claimed to know better than the workers what their interests were. With the exception of marginal mavericks, the anti-paternalist argument against workplace safety regulations and workers’ compensation is since long a historical phenomenon. (See Spurgin (2006) for a refutation of anti-paternalist argumentation for unsafe workplaces.)

Seat Belts

The seat belt was invented by the English engineer George Cayley (1773–1857), who intended it for use in airplanes. In aviation, the use of seat belts was uncontroversial from the very beginning. For instance, the very first aircraft of the US army, delivered by the Wright brothers in 1910, came with seat belts. From the

beginning of commercial aviation traffic, passenger seats were provided with seat belts. In the US, the first government regulation requiring passenger seat belts went into force in 1926. The use of seat belts has continued to be uncontested in air traffic. Flight attendants asking passengers to buckle up are seldom met with negative reactions. In spite of the standard phrase “for your own safety” used in the safety announcements, airlines do not seem to be accused of being instruments of the “nanny state.” In short, no one seems to demand the right to travel in an airplane without using a seat belt (Johannessen 1984; Vivoda and Eby 2011).

Already in 1885, seat belts could be found on some horse carriages. They were also used in racing cars as early as in 1922. However, the automobile industry did not equip their vehicles with belts. In the 1930s, physicians with experience of treating victims of automobile accidents started to advocate the mounting and use of seat belts in motor vehicles. However, not much happened until the 1950s, when designers developed safer and more convenient seat belts for cars. In the beginning, their constructions were largely based on the belts used in airplanes. In 1959, the Swedish engineer Nils Bohlin (1930–2002), who had a background in constructing ejection seats for airplanes, invented the three-point seat belt for cars. In the 1960s, more and more countries required all new cars to be equipped with seat belts, and car manufacturers started to include them as a standard (Johannessen 1984; Vivoda and Eby 2011).

In 1970, the Australian state Victoria adopted the first law making the use of seat belts mandatory (but only for drivers and front seat passengers). This was followed in the 1970s, 1980s, and 1990s by more and more countries. Today, the use of seat belts is mandatory in front seats in most countries, but many countries still do not require the use of belts in rear seats. In the beginning, considerable anti-paternalist resistance was mobilized against seat belt legislation. It was described as an infringement on liberty. For instance, one American opponent claimed that such legislation “violates the right to bodily privacy and self-control” of the drivers and passengers, and that it treated them in a “coercive, demeaning manner” (Solan 1986). The critics did not hesitate to oppose seat belt legislation for children. For instance, in a debate in the British parliament in 1988, the Conservative MP Gary Waller said that it “would be going too far” to require that parents refrain from driving more children in the car than there are belts for. He favored a noncompulsory approach and argued that “we should consider carefully whether a regulation is necessary or whether it would be better for advice to be given in the highway code, and for persuasion rather than compulsion to be exercised to ensure that children are carried in cars as safely as possible” (Waller 1988, cc. 591–598).

As seat belt laws have been introduced and enforced, this type of reaction has become increasingly uncommon. As Elvebakk (2015, p. 298) noted, opposition to obligatory seat belts is now “long forgotten in Europe.” In 2019, two ethicists summarized the situation as follows:

Within a few years, wearing seat belts became widely accepted and indeed endorsed in most countries. It became not only a legal, but also a social norm precisely because it was made compulsory and people started buckling up. . .

As often happens, people simply get used to and comply with new legal requirements even when they are initially opposed to them, and in the long run they see it simply as a social norm. (Giubilini and Savulescu 2019, pp. 237–238 and 243)

This appears to be an accurate description of the general picture, although there are still a few people, both in politics and in academia, who oppose mandatory seat belt legislation (Solomon 2010; Flanigan 2017).

Helmets

Helmets have been worn by soldiers for at least four and a half millennia (Gabriel 2007, pp. xiv and 80). There does not seem to have been any uproar or mutinies against orders to wear protective headgear on the battlefield.

In 1917, the US Navy provided workers in shipyards with hard hats to protect them from falling objects. In the early 1930s, some construction workers used homemade hard hats for the same purpose, whereas others used surplus WW1 helmets provided by their trade union. The Golden Gate Bridge in San Francisco, built in 1933 to 1937, seems to have been the first major construction workplace where the employer provided helmets to all workers and made their use mandatory. Advertisements from the 1930s onwards for helmets of different shapes and materials show that there was a demand for protective headgear (Snell 2018; Rosenberg and Levenstein 2010). Although some individual workers found hard hats uncomfortable, there does not seem to have been any ideological or organized resistance, and their use is now undisputed on construction sites and other workplaces with risks of head injuries. The requirement to wear them is seldom questioned, or as two researchers on occupational safety wrote:

Hard hats are different from other forms of P[ersonal] P[rotective] E[quipment]; people wear them even when they don't need to. There is no compliance issue with hard hats. They're cool. (Rosenberg and Levenstein 2010, p. 240)

The first motorcycle helmet seems to have been constructed by the English physician Eric Gardner. It was used for the first time in a race on the Isle of Man in 1914. After the race, Dr. Gardner received the following message from a colleague: "Every year the humdrum of medical practice in the Isle of Man is relieved by interesting concussion cases in our hospital from the Tourist Trophy Race; this year, thanks to your damned helmet, we have had none." (Gardner 1941) In the inter-war period, motorcycle helmets seem to have been used on racetracks but not much on roads.

In 1939, the Australian-British neurologist Hugh Cairns (1896–1952) started to work with brain-injured WW2 soldiers. Many of them were military motorcycle messengers. After studying their injuries, he emphatically proposed "for all motorcyclists, civilians and fighting Forces alike, the use of a crash helmet of the type worn by racing motor-cyclists" (Cairns 1941, p. 466). Following his advice, both the Army and the Air Force provided all their motorcyclists with helmets, which they

were required to use (Lanska 2009). This initiative seems to have been well received by the motorcyclists themselves. Cairns wrote:

When we began to treat Army motor-cyclists at [a military hospital in] Oxford, we naturally got to know their comrades—keen motor-cyclists in the Army Training Schools, Royal Corps of Signals, and other units, who were very much alive to the wastage of their man-power from accidents, and they needed little encouragement to become enthusiastic advocates of compulsory use of crash helmets. (Cairns 1946, p. 322)

Summing up his work after the war, Cairns wrote:

From these experiences there can be little doubt that adoption of a crash helmet as standard wear by all civilian motor-cyclists would result in considerable saving of life, working time, and the time of the hospitals. (Cairns 1946, p. 323)

After the war, motorcycle helmets were probably more common than before, and much improved helmets reached the market. However, a large number of motorcyclists continued to ride unprotected. This was changed through political decisions. During the period when seat belts became mandatory in country after country, motorcycle helmets also became mandatory (but often somewhat later than seat belts). Currently, motorcycle helmets are obligatory almost everywhere in the world, except in the US (<https://apps.who.int/gho/data/view.main.51427>. Accessed April 9, 2020).

In 1977, all but a handful states in the US had mandatory helmet laws that applied to all riders. However, there was considerable resistance from an active anti-helmet lobby, whose anti-paternalistic message appealed to Conservative and Libertarian lawmakers. In consequence, several states rescinded their helmet laws or weakened them to apply only to teenage drivers. In 1995, Congress repealed the federal legislation that had incentivized states to uphold the helmet requirement. This led to an avalanche of revocations of state laws (Jones and Bayer 2007). At the time of writing (2020), only 19 states and the District of Columbia have universal helmet laws (<https://www.iihs.org/topics/motorcycles>. Accessed April 9, 2020).

Just like motorcycle helmets, bicycle helmets protect efficiently against severe brain damage and fatalities that may result from cycling accidents. In bicycle sport, a minority of riders have used helmets at least since the 1950s, and helmets became much more common in the 1990s. They became obligatory in 2003 (https://web.archive.org/web/20160304110024/http://oldsite.uci.ch/english/news/news_2002/20030502i_comm.htm. Accessed April 9, 2020.). In the 1970s, a small minority of non-racing (transportational and recreational) cyclists began to use helmets, and a few companies sold helmets intended specifically for cycling. Standards were adopted, and helmets satisfying the standards were designed. In Victoria (Australia), helmet promotion activities led up to the introduction of a law in 1990 that made helmets mandatory for all cyclists (Cameron et al. 1994). The other Australian states have followed suit, and universal helmet laws have also been adopted in New Zealand and Argentina. Some other countries have made helmets

mandatory only for children, but still in 2020, most countries have no legislation requiring the use of bicycle helmets.

Ice hockey has been played since the late nineteenth century. It is a sport with many head injuries, but up till and through the 1950s, almost no player ever used a helmet or any other type of head or face protection. This was much due to a “culture of toughness” that led players to view injuries as a part of the sport, which one had to endure. In the 1960s, the major hockey organizations introduced obligatory helmets for youth players. In 1979, the National Hockey League made helmets obligatory for professional players (with an exemption for older players, who were allowed to continue playing without a helmet). Helmets are now an uncontroversial part of the equipment of hockey players, and there is no sign that players crave for playing without them (Bachynski 2012).

Drunk Driving

Patricia Waller has described the attitudes to drunk driving that still prevailed in many countries in the 1970s:

Drunk driving was considered more or less a “folk crime,” almost a rite of passage for young males. Most adults in the United States used alcohol, and most of them, at some point, drove after doing so. This is not to say that they drove drunk, but many of them undoubtedly drove when they were somewhat impaired. Although the law provided for fairly harsh penalties, they were rarely applied. Upon arraignment, defendants would ask for a jury trial, and because drinking and driving was so widespread, juries almost invariably acquitted the defendant, thinking, “There but for the grace of God go I.” (Waller 2001, p. 3)

Writing in 2001, she noted that “[r]emarkable progress has been made” (ibid., p. 4). This would not have been possible without the convincing data produced by the research community, she said, but another contribution was at least as important:

It was citizen action groups that provided the impetus for major changes in public policy governing drinking and driving. Their activities generated public support for enforcement of existing laws and enactment of new ones. (Waller 2001, p. 3)

Particularly in the 1980s, activist groups changed the focus of the discussion from the drivers to the victims, not least children (Williams 2006). As emphasized by Leonard Evans (1991, p. 352), “the most important factor was the widespread serious discussion of the tragic dimension of the problem in the media, with the consequent deglamorizing of the drunk as a likeable humorous character.”

A few contrary voices have been heard. In 1994, a philosopher published a journal article in which he claimed that drunk driving is not a serious offense. As an alternative to preventing drunk driving he proposed the removal of roadside trees, which are the “most commonly struck objects” by drunk drivers, as well as the introduction of airbags, which “would make drunk driving much less risky” (Husak 1994, p. 68). As late as in 2011, a major libertarian magazine in the US published an article describing

the prohibition of drunk driving as “nanny state” legislation. The author claimed, incorrectly, that “[e]xperienced drinkers” differ from less experienced ones in that they “can function relatively normally with a B[lood] A[lcohol] C[ontent] at or above the legal threshold.” He used this as an argument for entirely “repealing drunk driving laws” so that a drunk person should be allowed to drive unless he is “violating road rules or causing an accident” (Balko 2011). However, in experiments with driving simulators, heavy (“experienced”) drinkers did not perform better under the influence of alcohol than unexperienced drinkers (Marczinski and Fillmore 2009; Fell and Voas 2014).

Today it is difficult to defend drunk driving in public. In some countries, opponents of efficiently enforced drunk driving laws have instead turned against the means of enforcement, for instance, claiming that random sobriety checks on the roads (sobriety checkpoints) are an unbearable infringement of the liberty of drivers. In Australia, activists working for the introduction of random breath testing found themselves opposed by “the alcohol industry and its supporters, notably civil libertarians and anti-nanny state proponents,” who claimed that this would be an infringement on individual freedom (Howat et al. 1992, p. 20). One American legal scholar described this method of law enforcement as “[f]ascist-like sobriety checkpoints” (Miller 1993, p. 174). This argumentation has not been without success; in several states in the US the police are not allowed to perform random sobriety checks (Fell et al. 2004; Vissing 2014).

In the 1990s, a reliable method to prevent drunk driving was introduced in the form of alcohol interlocks (breath alcohol ignition interlock devices). Currently available technology requires the driver to provide a breath sample prior to starting the vehicle (and usually also at random times when driving, to prevent the use of another person’s exhalation air). Ongoing research aims at obtaining the same result without a breath sample – either through analysis of the driver’s normal exhalations or by measuring blood alcohol levels noninvasively with infrared light directed at the driver’s fingertips (<https://www.dadss.org>. Accessed 10 April 2020.). In a large number of countries as well as all states in the US, courts can order convicted drunken drivers to install alcohol interlocks in their cars, as a means to prevent repeat offence. In 2005, the Swedish (social democratic) government announced a coming bill that would make alcohol interlocks mandatory in new cars by 2012. However, the new (conservative) government that entered office after the general election in 2006 decided not to go forward with the legislation, and since then there has not been a political majority for the measure (Grill and Nihlén Fahlquist 2012).

As yet, there does not seem to be much resistance against alcohol interlocks. That may be because their current use has the effect of increasing the automotive freedom of convicted drunk drivers. By installing the alcolock, they can keep or regain their driver’s license, which they would otherwise have lost. However, it is plausible that proposals to introduce breathalyzers in all motor vehicles can give rise to counter-reactions of the same types that we have seen against laws mandating seat belts and motorcycle helmets. In 2019, the European Union announced plans to require all new cars from 2022 to be technically prepared for easy installation of alcohol interlocks. This will facilitate later installation of interlocks if the driver is convicted

of drunk driving, or chooses to mount an interlock for other reasons (such as, possibly, a future tax reduction for those who do so). It will of course also simplify any future decision to make alcohol interlocks obligatory. This latter possibility was keenly observed by a conservative motorist organization in the US, the National Motorists Association (NMA). In one of their newsletters, they first criticized other safety features mandated in Europe, such as advanced emergency braking systems (AEBS) and systems that warn a driver who gets drowsy. They continued:

The other mandatory element that is problematic: Every new car would include an alcohol interlock installation facilitation. The driver's blood alcohol content could potentially be checked by breathalyzer or by tactile sensors. The mandate would mean no driver could start a car without passing the electronic testing... Vision Zero proponents want to rein in irresponsible, that is, all (in their minds) motorists by literally limiting speed, choice, privacy and personal responsibility. The war on cars and on drivers is heating up. (National Motorists Association 2019)

Speeding

It was known long before the introduction of the modern automobile that high speed increases the risk of accidents involving vehicles. The oldest speed limit on record seems to be a regulation that was passed in Newport, Rhode Island, in 1678, forbidding horse-riders to “gallop or to run speed” in the streets. In Boston in 1757, carriages were limited to “foot pace” on Sundays to protect church visitors. The first speed limit referring to an exact velocity was probably the legislation passed in Britain in 1865 for steam carriages on highroads, which were limited to at most 4 mph (6 km/h). In the early twentieth century, many countries introduced speed limits for automobiles, often with different velocities for different kinds of roads (Elston 1971, pp. 21–22). This was “a time when motorists were in the minority and unpopular” (Tripp 1928, p. 535), and judging by articles in the newspapers at the time, there was considerable public support for these restrictions. In 1907, the American magazine *Harper's Weekly* published an article about “speed mania,” describing the perpetrators as follows:

[I]n every case their mania arises from an overweening sense of their own importance, accompanied by very slight capacity for self-restraint. The type of man who motors at dangerous speed is the same type that speculates in more stocks than he is able to carry, eats and drinks more than he can assimilate, covers himself with gaudy jewels, makes an objectionable exhibition of himself on every possible occasion. The strong arm of the law is the only effective curb for this species... (Underwood 1907)

In contrast, “[a] decent regard for the safety of mankind will always preserve the normal man from giving way to speed mania” (ibid.). It did not take long before motorists started to strike back. In 1911, the Automobile Club of America demanded that “all speed limitations should be omitted from the law,” for the somewhat curious reason that “no matter what the limit is, the majority of drivers will try to exceed it.”

They did however emphasize that “[c]areful driving” should be enforced (Dorrian 1911).

Speeding continues to be a major factor contributing to the prevalence and severity of road traffic accidents (Farmer 2017; de Bellis et al. 2018). Cavalier attitudes to the dangers of speeding are perpetuated by widespread and entrenched myths such as “a general belief that speeding slightly in excess of the limit (up to at least 5 km/h, perhaps as much as 10 km/h) is not associated with increased crash risk if the driver is otherwise driving safely” (Delaney et al. 2005b, p. 23). A Swedish writer who obviously believed in this myth claimed that speed limits are a form of “collective punishment” in order to solve the problem with “those who do not manage reasonably safe” at higher speeds (Eberhard 2006, p. 291).

Speed cameras, which have been introduced throughout the world since the 1970s, have met with considerable resistance. The experiences reported from an Australian speed camera program appear to be typical:

As the speed camera program expanded in 2001 and lower speeds over the limit were targeted, various controversies came to the fore and had high public profile. These included the notions that the program was established principally to raise revenue for the government (fine money goes to consolidated government revenue), that camera locations included those where it was “safe” to speed, that overt operation of cameras was most appropriate if the aim was to deter speeders at unsafe locations, that exceeding the speed limit by only a small amount was safe, that there was no opportunity afforded to explain the circumstances of the event, and that the reliability of cameras and speedometers came into doubt when the tolerance was perceived to be only about 3 km/h (1.9 mi/h) above the speed limit. (Delaney et al. 2005a, p. 408)

This is a long list of complaints, but interestingly, it does not contain any clear anti-paternalist or “nanny state” argumentation. In Britain around the year 2000, resistance to speed cameras was even more outspoken and active than in Australia, but the argumentation was about the same (Pilkington 2003). Even the clandestine organization Motorists Against Detection (MAD), which claimed to have destroyed 600 speed cameras, justified their activities by claiming that the real purpose of the cameras was to collect speeding fines for the government (Chancellor 2003). Speed cameras are still vandalized in many countries (Max 2019; Robinson 2019). In summary, there are people who hate speed cameras so intensely that they go to the extreme of vandalizing them, but they do not usually invoke anti-paternalist arguments to justify their opinions or their actions.

Speed control can be taken to a new level with Intelligent Speed Adaptation (ISA), i.e., a system that keeps track of the speed limit where the car is driven and relates it to the actual speed. ISA systems can be either passive or active. A passive system warns the driver, whereas an active system reduces the speed to the legal limit (possibly with an option for the driver to override the system). Active ISA systems (speed limiters) have been predicted to drastically reduce the number of fatalities in road traffic (Carsten and Tateb 2005). Based on this, a strong moral argument can be made in favor of making active ISA obligatory (Hansson 2014, p. 373; Smids 2018). However, several studies indicate that drivers’ willingness to install and use such a

system is low (Fu et al. 2020). Numerous newspaper articles describe speed limiters as an entrance gate to the “nanny state” (e.g., Mowat 2019). A Canadian truck-driver who had been ordered to use a speed limiter went to court, arguing that this order infringed on his freedom (Nyholm and Smids 2020). Although he lost his case, this indicates that the introduction of automatic speed adaptation can lead to a clash between traffic safety objectives and strongly felt anti-paternalistic sentiments.

Conclusions

Summing up the above, we have studied quite a few examples of legal requirements and mandatory practices that can be said to reduce the freedom of choice of individuals in some way or other. We can classify them in three main categories. The first category consists of *uncontested restrictions*, i.e., restrictions in liberty or freedom of choice that have never encountered significant ideological or organized opposition:

- helmets in the military
- seat belts in airplanes
- helmets in motorcycle racing
- helmets on construction workplaces
- alcohol interlocks for convicted drunk drivers

These examples have somewhat different backgrounds. Soldiers have been ordered to wear helmets since ancient times, and no one seems have come up with the idea of accusing their commanders of thereby restricting the soldiers’ liberty, or freedom of choice.

Seat belts were mounted on passenger airplanes from the beginning of commercial aviation (at a time when flights were much more shaky than today), and today their use is an entrenched habit, which no one seems to put in question. They are just one of a large number of safety arrangements that we take for granted without thinking much about them. To see how common and how generally accepted such arrangements are, let us suppose that you try to avoid them. Suppose that you go to a household appliances store and look for a microwave oven that can emit microwaves when the lid is open. You will be told that no such ovens are made. If you go to the hardware store and look for power tools with uncovered live electrical parts, you will be equally disappointed. All the power tools on sale have sturdy protective shields to keep the user safe from electric shocks. Next, you go to an auto glass shop in order to have the cracked windshield of your car replaced. You want a windshield of common inexpensive window glass, instead of the expensive laminated safety glass that you are offered. “I am sorry,” says the repairman, “but all windshields are made of laminated glass.”

What has happened? Has the nanny state invaded the shops and stopped them from providing the articles you asked for? Yes, in a sense it has, since there are (in most countries) regulations enforcing safety standards for microwave ovens,

power tools, windshields, and a myriad of other products. However, even in the absence of such regulations, it is much to be doubted whether any of the articles you asked for would be manufactured, or whether, if they were made, anyone would buy them. No one seems to care for the liberty to avail themselves of such products. And similarly, no one seems to ask for the liberty of remaining unbelted in an aircraft during take-off or landing. In fact, few would even care to know if it is the government, the airline, or perhaps some international aviation organization that decided that we all have to use a safety belt onboard.

The last three examples in this group differ from airplane seat belts in having been introduced as a new practice in an existing activity. Motorcycle racing and construction work were originally performed without helmets. The protective headgear was introduced without any ideological or organized opposition. Similarly, in-car breathalizers for convicted drunk drivers were introduced without resistance.

Our second major category consists of *previously contested* restrictions in liberties and freedom of choice that are now generally accepted:

- clean water and hygienic treatment of sewage
- employer's responsibility for safe workplaces
- seat belts in cars (in most countries)
- helmets in ice hockey
- prohibition of drunk driving

There have been days when fervent campaigners fought for the liberty to forgo modern sanitation, to work in extremely dangerous mines (or rather, to hire others to do so), to drive drunk, and to play ice hockey without helmets. Today, no one (with the possible exception of a few eccentrics) propagates these standpoints, for the simple reason that desires to exercise these liberties are no longer considered to be reasonable or worth supporting. The same is now, in most countries, true of desires to travel without seat belts in motorcars. This means, for instance, that hockey helmets are now equally uncontroversial on the ice rink as hard hats on the construction site. Notably, in all these cases, the transition has been gradual, and it has been accomplished by activists concerned about other people's health and safety.

Our third and final category contains the *currently contested* restrictions in liberties and freedom of choice. These are the practices that encounter significant ideological or organized resistance:

- motorcycle helmets on roads
- bicycle helmets
- universal alcohol interlocks
- speed cameras
- speed limiters in cars

Judging by the historical experience, some of these practices may in the future find their way into the second category. For instance, in countries such as Sweden, motorcycle helmets seem to be at a rather late stage in that process. However, this

is not a transition that can be taken for granted. Vaccination belongs to this third category. It has remained contested for more than two centuries, despite its tremendous contributions to public health (Meyer and Reiter 2004).

At least four important conclusions can be drawn from this résumé of uncontested, previously contested, and currently contested infringements on liberty. First, *some but far from all restrictions of individual liberties are at all contested, and over time, there are large changes in what restrictions are contested*. We saw this clearly in the discussion of the second category (previously contested).

Secondly, *which restrictions are contested has very little to do with the distinction between paternalistic and non-paternalistic constraints on liberty*. Some of the most passionately opposed restrictions cannot with any vestige of credibility be described as paternalistic. This applies, for instance, to prohibition of drunk driving, as well as speed cameras and speed limiters. On the other hand, some of the restrictions that have encountered little or no opposition are commonly considered to be paternalistic. This includes seat belts in airplanes and hard hats in the building industry.

Thirdly, *we accept restrictions more easily if we encounter them in situations where we are accustomed to do as we are instructed*. This goes a long way towards explaining why helmets have been accepted in armies and on workplaces, as well as – after some initial opposition – in organized sport activities, where directives from referees and coaches have to be followed (Bachynski 2012).

Fourthly and finally, *we are less prone to accept restrictions if they require changes in our entrenched habits*. This seems to have been an important factor in resistance to several safety measures in road traffic, such as seat belts, helmets, sobriety, and reduced speed. As a corollary, we can expect restrictions to encounter little opposition if they are introduced from the beginning in a new activity (e.g., seat belts in airplanes). Habits also seem to have a role when previously disputed restrictions gain acceptance. For instance, seat belt laws induced law-biding people to buckle up. This became a habit – the “new normal” – and habituation to this convenient and well-justified routine appears to have led to its acceptance (Boughton 1984, p. 187).

Seen in this perspective, many of the issues that have been depicted as conflicts between paternalism and anti-paternalism are in fact largely fought along other – no less important – conflict lines, such as that between one person’s freedom and other persons’ safety. However, this does not imply that the issue of paternalism is irrelevant. In the next section we are going to have a close look at the scope of anti-paternalist arguments.

The Significance of Human Connectedness

The purpose of anti-paternalism is to protect everyone’s right to make and follow her own choices in matters that only concern herself and nobody else. Therefore, the scope and importance of anti-paternalism depends on what types of situations there are that answer to that description.

Obviously, there are many choices that only concern the person who makes them. This includes many of the liberties we have in private life. The color of my coat, what books I read, what furniture I buy (if I can afford it and I live alone) are under normal circumstances nobody else's business, and these are only a few of the many liberties that we should all have. There are also many liberties that we should all have, although our exercise of them can have large impact on others; this includes the political liberties of democracy. However, many of the allegedly private actions that have been at focus in discussions on paternalism are not as unconnected to other people's lives as they have often been claimed to be. This, as we will see, can have considerable impact on the scope and strength of anti-paternalist arguments.

Some discussions of liberty give the impression that we humans are rational atoms, each with her own, individually chosen, goals, and with no other obligations to each other than to avoid colliding when we follow our chosen trajectories in the huge space of possibilities. But this is not a true picture of human life. To the contrary, we are all bound to each other in innumerable ways, and no one can reach very far without the – past or present, willing or unwilling, intended or unintended – actions by others. Sometimes this interdependence is hidden from us, and sometimes we prefer to see it as an unchangeable background, like the mountains and the seas. But that is an illusion. Almost everything we do connects us with other human beings.

This interdependence is reciprocal, which means that it goes in two directions. On the one hand, the choices and decisions that we make are far less individual than what we tend to believe, and much more co-determined by others. On the other hand, our actions have effects on others, effects that we are often not aware of. Both these aspects of our social interdependence have important moral implications, not least for the scope and applicability of anti-paternalist arguments. Let us begin with the first of them.

Combined Causes and Extended Anti-paternalism

The causality of human actions is much more complex than what we usually think (► Chap. 5, “Responsibility in Road Traffic”). Most of the actions that we ascribe to a single person do in fact depend on the combined actions of many persons (sometimes acting at different points in time). In a discussion on paternalism, we have reason to apply this insight in particular to self-harming actions, or more generally, actions that go against the interests of the person who performs them. (There are of course many cases in which it is impossible to determine univocally what goes against a person's interests and what does not. For our present purpose, we can leave such cases aside.) In many cases when we talk of a person as harming herself, or acting against her own interests, others have made choices and taken actions without which the self-harm would not have been possible:

- *A drug addict harms herself by buying addictive drugs that destroys her health.*

This would not have been possible without drug dealers who choose to sell products that havoc the health and the lives of most of the people who use them.

- *A smoker harms her own health by smoking cigarettes.*

This would not have been possible without the cigarette companies that have chosen to sell and aggressively market products that kill about half of the people who use them.

- *A cyclist rides to work on a dangerous road among cars and trucks driving at high speed.*

This would not have happened if the road authorities had arranged a separate bicycle lane.

- *A motorist drives without a seat belt.*

In the vast majority of cases, she would have fastened the seat belt if the manufacturer had equipped the car with a seat belt reminder.

- *A motorist drives at an excessively high speed on the expressway, endangering both her own life and that of others.*

This would not have happened if the manufacturer had installed a speed limiter making it impossible to drive the car at speeds at which it cannot be controlled.

These are all cases of “mixed” causality, in which harm to a person is caused by a combination of her own actions and actions by others. In such cases, the causal role of the others is often downplayed, and the action is described as purely self-harming. For instance, the drug dealer is sure to claim that “it was her own decision to buy the drug” and “if I hadn’t sold it to her, then someone else would have done so.” (For a discussion of responsibility ascriptions in such cases, see Hansson (► Chap. 5, “Responsibility in Road Traffic”).)

Anti-paternalism, properly speaking, defends a person’s right to harm herself. From such a right it does not follow that others have a right to harm her, or to contribute to harming her. However, attempts have often been made to extend the “green light of liberty” that anti-paternalism bestows on self-harming actions to other-harming actions by others that contribute to the same harm. One of the most remarkable examples of this is the argumentation of the tobacco industry. This is an industry with a long and still on-going record of aggressively marketing massively death-bringing products. For instance, in three decades, the 1970s to 1990s, tobacco companies focused disproportionately large resources on areas in American big cities with predominantly African American residents. They distributed free cigarettes in the streets, and focused particularly on what they called the “starter market,” i.e., new smokers, mostly minors. These marketing activities led to a considerable increase in smoking among Afroamericans. One of the consequences of this was that the prevalence of lung cancer among black men, which had previously been lower than among white men, rose to levels higher than those among white men (Yerger et al. 2007). This is only one of many examples of the decisive impact that marketing and propaganda has on smoking. However, the tobacco industry spends considerable efforts on describing smoking as depending exclusively on the smokers’ own free decisions, claiming, for instance, that “the growing intrusion of government in the lives of adult smokers is a threat to the freedoms of all citizens” (Katz 2005, p. ii33. Cf. Cardador et al. 1995; Apollonio and Bero 2007; Schneider and Glantz 2008; Fallin et al. 2014). With these campaigns they try to achieve two results: First, they

try to attribute smoking entirely to the individual smoker's decision. They have been remarkably successful in promoting this rather obvious fallacy. For instance, in a book published in 2006, a Swedish physician (!) ridiculed lung cancer patients who sued tobacco companies, claiming that these plaintiffs professed "not to know that it was dangerous to smoke" (Eberhard 2006, p. 313). Secondly, by appealing to anti-paternalism on behalf of the smokers, they try to ensure that anti-paternalism also protects their own, massively other-harming, activities.

This is an extreme example, but the same pattern of thought can be found in many other contexts where anti-paternalism is invoked to protect other-harming actions. But how valid is this type of argument? If anti-paternalism protects my right to drink so much that I am unable to walk home, does it also absolve the publican who serves me that much liquor? If anti-paternalism allows me to eat food with toxic mold or pesticide residues, does it also acquit the food industry or my grocer for selling such food (properly labeled) to me? If anti-paternalism lets me ride a motor-cycle without a helmet, does it also exonerate the rental agency that rents out a motorbike to me without including a helmet in the rental? And if it permits me to drive a car without a seat belt, does it also condone a carmaker who refrains from mounting a seat belt reminder as standard equipment?

The general issue illustrated in these examples is that of *extended anti-paternalism*, by which is meant the use of anti-paternalist arguments to justify actions or activities that harm (consenting) others, usually in combination with self-harming actions or decisions by the persons in question. The term "extended anti-paternalism" was introduced in Hansson (2005). In the literature on paternalism, the distinction between anti-paternalism in the proper sense and this extension has usually not been made. Authors who observed the distinction, or some variant of it, have usually not made much of it (Mill (1977 [1859], pp. 296–299; Dworkin 1972, pp. 67–68; Dworkin 2020, Sect. 2.4; Feinberg 1986, pp. 9–11; Schramme 2015).

From a logical point of view, it is obvious that a right to harm oneself does not imply a right to harm consenting others or to facilitate or contribute to their self-harming activities. But logic alone cannot tell us if it is reasonable to extend anti-paternalism in this way. That is a moral question, and it has to be discussed in terms of moral principles. We can approach this moral issue by first reminding ourselves of how we deal with it in private life. What moral attitudes do we usually have to the combined effects on a person of actions by herself and actions performed by others? It does not take much reflection to see that this depends on the nature of the combined effect.

We are at liberty to do many things that are presumed to have a positive value. For instance, suppose that you have a friend who likes to read books. We would normally see it as not only allowed, but also commendable, to help her to exert that liberty, for instance, by picking up books at the library for her or recommending her books that she might like. We also see it as positive if (commercial or nonprofit) services, such as bookshops, libraries, and reading groups, are available for her and others wishing to exert this liberty. The same applies to a multitude of other activities that people are at liberty to perform, and which we presume to have predominantly positive effects on the well-being of those who choose them. This is how we usually perceive actions that help people to exert their liberties.

But there are also liberties that are endorsed *per se* (for the sake of liberty as such), rather than for any positive effects of realizing them. Liberties to harm oneself belong to this category. In private life, helping others to harm themselves is usually held to be morally misguided, if not outright wrong. Suppose that your neighbor has the unusual obsession of burning small scars on his body with a burning glass. This action is presumably not illegal, and neither do we tend to morally condemn it. (On what grounds could we condemn it, and of what use would it be to do so?) Let us assume that we take an anti-paternalist position to this activity; in other words we consider it to be a form of self-harm that he is at liberty to perform, and which we have no right to prevent him from. Is it then reasonable to augment our anti-paternalism in this case to extended anti-paternalism? In other words: Is it morally allowable to help him by holding the burning glass over his back so that he can get scars there as well? Is it OK to recommend him a stronger burning-glass, or buy one for him? And if he becomes an Internet celebrity and attracts followers burning scars into their own skins, is it good business ethics to produce and market “scar burners,” specially adapted to the purpose?

The obvious answer to all these questions is “no.” In private life, morality requires that we care for other people, and doing so is not compatible with encouraging or helping them to harm themselves, or even with looking the other way when they do so. The morally laudable reaction to his burning obsession would be to try to help him: talk to him about why he does it and what could to make him stop, find ways to help him break the habit, perhaps find a psychiatrist who can help him.

The example is unusual, but it illustrates how we usually react to clear examples of self-harming activities when we encounter them in our private lives: Instead of contributing to other people’s self-harm, we try to help them out of it. We do not praise those who buy liquor to the alcoholic, emetics to the bulimic, or poison to the depressed and suicidal friend. Instead, we praise those who try to make them change their minds and develop more positive thoughts and habits. Extended anti-paternalism does not work in our everyday lives. Empathy and care do.

This insight can also be applied on a larger social scale, in discussions about what communities, organizations, companies, and governments should do. By breaking up the false inference from anti-paternalism to extended anti-paternalism, we open up for a more efficient strategy to promote safety and public health. Such a strategy should actively prevent the exploitation of self-harming action. This includes, for instance, effective curtailment of the pathogenic activities of tobacco companies. Instead of blaming each other for dangerous and self-destructive behavior, we should make our society and the options it offers less dangerous. This is the approach that has been applied for more than a century in occupational safety, with considerable success. Instead of admonishing workers not to put their fingers into the press machine, the employer installs a machine that does not move unless both hands are in safe places. If this approach was ever criticized for being paternalistic, then that criticism is long since forgotten. Applying it in road traffic would mean to ensure that vehicles and the road system are so constructed that self- and other-harming activities on the road, such as speeding, driving drunk or on the wrong lane, are technically impossible. This, by the way, is a large part of what Vision Zero is all about.

Herd Effects: How We All Influence Each Other

In the discussion on paternalism in public health, it has often been pointed out that many activities that seem to be “pure self-harm” with no effect on others do in fact have significant impact on others than the person who exposes herself to a harm or risk. An unbelted back seat passenger can become a projectile in the event of a crash, and injure or even kill front seat occupants. This effect is surprisingly large; according to one study, the risk for a belted driver or front seat passenger to die in a crash increases almost fivefold if she has an unbelted passenger behind her on a rear seat (Ichikawa et al. 2002). Psychological costs to bystanders who witness a deadly crash have also been mentioned, and should most certainly be taken seriously (Feinberg 1986, pp. 139–141). An even more serious concern is the suffering of family members: children who lose a parent and spouses who have to care for a seriously injured accident victim, or mourn the loss of their partner (McKenna 2007). The risk of making one’s children father- or motherless is probably an important consideration in many private discussions and deliberations. It was also discussed in the newspapers after the death of a female climber, the mother of two small children, in a descent from the summit of K2 in 1995 (Barnard 2002). However, it does not seem to have been much discussed in connection with accidents in road traffic such as the around 2000 unhelmeted motorcyclists who die on American roads every year, many of whom are sure to be parents (NHTSA 2019, p. 11). Arguments referring to the economic costs have had a more prominent role. They concern the costs of hospital care, rehabilitation, a future life in a nursing home, and subsistence for family members (Hundley et al. 2004). In some of the American states that allow unhelmeted motorcycling, unhelmeted riders are required to have an extra insurance to cover personal injuries (Langland-Orban and Flint 2011). Obviously, if these psychological, social, and economic costs of injuries and fatalities on the road are taken into account, then the common anti-paternalist argumentation against measures to improve traffic safety will lose almost all its bite.

However, arguably the most important reason why anti-paternalism cannot be applied to actions such as helmetless motorcycling has been absent from these discussions, namely, what I propose to call the *herd effects* of individual behavior. Many authors have provided what appears to be intended as a complete list of effects of helmetless motorcycling that anti-paternalism cannot legitimize, without mentioning herd effects (Schonsheck 1994, pp. 107–141; Camerer et al. 2003; de Marneffe 2006, pp. 82–83; Thaler and Sunstein 2008, pp. 232–233; Nolte et al. 2017).

Herd effects arise because in almost all our decisions, we humans tend to follow the examples of others. This has been known for long. Already Julian of Eclanum (c.386–c.455) wrote about the “contagion of sin” (Barclift 1991, p. 14), and his contemporary John Chrysostom (c.349–407) strongly emphasized the positive side of this, namely, the importance of being a good example to others (Chrysostomus 1836, p. 788, c.1D). The concept of behavioral or social contagion was introduced into social science by the French scientist Gustave Le Bon (1841–1931). Today there is an extensive literature on how all kinds of human behavior is propagated in

populations through our tendency to copy or imitate the behavior of others (Lehmann and Ahn 2018).

This applies not least to behaviors relevant to public health, such as smoking, drinking, substance abuse, and physical exercise (Christakis and Fowler 2013). There is also ample evidence of strong contagion effects on both the following and the breach of safety rules on workplaces (Liang et al. 2018; Liang and Zhang 2019). In road traffic, social contagion has a large role in creating a “culture of driving.” An important study in the early 1990s showed that “drivers are sensitive to the influence of others and... that a small shift in the behavior of few can be amplified, through the interaction between individuals and their collectives, to a larger effect, resulting in a changed social environment or a modified ‘culture of driving’” (Zaidel 1992, p. 585). Other studies have confirmed that drivers have a strong tendency to adapt to the prevailing speed pattern (Connolly and Åberg 1993; Edwards et al. 2014). There is also evidence of similar effects on risky pedestrian behavior McGhie et al. 2012).

Unusually clear evidence of social contagion has been found in those states in the US that replaced a law requiring all motorcyclists to carry a helmet by a law that only required this of young riders. In these states, large numbers of young riders chose not to use a helmet when they saw older riders riding helmetless. In Florida, the number of motorcycle fatalities among riders younger than 21 nearly tripled after the helmet law was downgraded in this way in 2000 (Bachynski 2012, p. 2216. Cf. Nolte et al. 2017). A group of public health researchers have pointed out that a similar effect can be expected for bicycle helmets:

We also conjecture that by applying the [bicycle helmet] law to children but not adults we would encourage a “rite of passage” effect (much as happens with cigarette smoking), whereby older children abandoned helmets to signify their maturity. This perverse effect would subvert an explicit, if secondary, policy aim of making helmets compulsory for children, which is to encourage adults to adopt helmets in order to set an example without compromise to adults’ legal liberty and moral autonomy. (Sheikh et al. 2004, p. 264)

These results put what seemed to be purely self-affecting decisions, such as whether or not to use a seat belt or a helmet, in an important, previously neglected, perspective, namely, that of how our actions influence other people’s actions. We humans are not mental solitarians who make decisions and choices independently of each other. To the contrary, we are herd creatures. We all observe what others do, and we usually follow trends, in particular trends among people whom we feel akin to. If you speed, others will feel encouraged to do the same. If you use a bicycle helmet, then that can have a considerable influence on people whom you know, and it can also have a smaller influence on a larger number of people, namely, those who see you riding with the helmet.

Do these herd effects have any moral relevance? Do we have a moral obligation to behave in ways that give rise to positive rather than negative effects on our fellow beings? Or are these effects just byproducts of our actions and choices, for which we have no responsibility? I can see no reason to treat herd effects differently than other consequences of human action. In other words, they are morally relevant. This means that we are morally accountable for the foreseeable effects of the examples

we set. It is morally blameworthy to set bad examples, for instance, by taking foolish risks to no avail, and it is morally praiseworthy to set good examples, for instance, by following safety rules and using protective equipment. This is not a new or original standpoint; presumably it has never been a highly esteemed conduct to drink oneself into a stupor or play dangerously with knives in the sight of children or adolescents. But apparently we need to be reminded that the same moral principles apply to many behaviors that have been considered protected by anti-paternalism.

A useful comparison can be made with another type of contagion, namely, infectious contagion. There are two reasons why you should take vaccinations against infections that threaten public health. First, vaccination protects your own health. Secondly, vaccination protects other people, since if you do not get the disease you will not spread it to others. This is particularly important for vulnerable people who cannot, for medical reasons, take the vaccine. For instance, vaccination against measles, a potentially deadly childhood disease, cannot be given to babies. Their only protection is herd immunity in the population, which is obtained when those who can take the vaccine do so. This means that vaccination is not just a personal matter, and consequently refusal to vaccinate is by no means protected by anti-paternalism. There are strong reasons to regard vaccination against diseases that threaten the community as a moral obligation (Jamrozik et al. 2016; Pierik 2018).

This is written during the Covid-19 pandemic in 2020, which has led to drastic measures all over the planet to contain the infection, including lockdowns and social distancing. The same reasons for these measures apply as for vaccination: You do this in order to protect yourself against the disease, but also to avoid spreading it to others. In this case, vulnerable, mostly elderly, people are most at risk, and they depend for their protection on the willingness of others to take the measures necessary to avoid contracting and spreading the disease.

Unfortunately, there is in both cases, vaccination and social distancing, a minority of people who either have not understood that the recommended measures are needed to protect others than themselves, or just do not care. We can call this the *Bolsonaro fallacy* after the Brazilian president Jair Bolsonaro who has flagrantly and defiantly violated the directives of his own administration to reduce the spread of Covid-19, with the justification “If I have got myself infected, so OK? Look, that is my responsibility, no one has anything to do with this.” (“Se eu me contaminei, tá certo? Olha, isso é responsabilidade minha, ninguém tem nada a ver com isso.” From an interview on March 16, 2020. Last accessed 14 April 2020 on <https://www.youtube.com/watch?v=M0za8MSO64&feature=youtu.be>. At 6:58–7:02.)

This example provides a useful background for the moral appraisal of how we contribute to social contagion, which is no less a formidable force in society. The mechanisms through which our behaviors impact on those of others can be even more imperceptible than viruses and bacteria, but they are just as real. A motorcyclist who wears a helmet contributes to creating a culture, or social environment, in which helmets are an unquestioned part of motorcycling, just as hard hats are worn routinely on construction sites. A motorcyclist who refrains from using a helmet contributes, to the contrary, to maintaining or creating a social environment in which it is normal, perhaps even “tough” or “cool,” to ride unhelmeted. The analogy with

taking vaccines and complying with social distancing in times of an epidemic is obvious. The vaccine shot that you receive reduces the risk of contracting the disease by a small amount for each of the persons who could potentially be infected by you, and the sum of all these small amounts matters morally. Similarly, the helmet on your own head encourages others to do the same, thereby contributing to making wearing a helmet the normal and expected thing to do. Social contagion is just as real, just as unavoidable, and just as morally important, as infectious contagion.

For both social and infectious contagion, what matters is the sum of a large number of small effects. But in both cases, we can see from the sums that these effects are real. Countries with high vaccination rates are almost unaffected by diseases that have devastating effects in unprotected populations. And in countries like Sweden, where young people may never have seen an unhelmeted motorcyclist, the question of riding helmetless does not even arise. In neither case does the fact that each of the many small individual effects cannot be perceived make them morally irrelevant (Hansson 1999). Therefore, a motorcyclist who refrains from wearing a helmet, claiming that this concerns no one but herself, makes the same mistake as the vaccine refuser and the organizer of crowds in a pandemic, namely, the Bolsonaro fallacy.

Putting this more positively, the fact that we are so much influenced by each other is a reason to put more weight on the old but sometimes forgotten moral value – and moral imperative – of setting good rather than bad examples. Asking a person to use a helmet, or to fasten her seat belt, is not a meddlesome intrusion into her private life. It is an invitation to take part in a joint effort to create or maintain a custom that protects us all. It is for obvious reasons better if such customs can be achieved and preserved without the force of law. This is how ice hockey helmets were introduced. However, if this does not work, then we should not forget what we have governments and laws for. We certainly need them when children contract potentially life-threatening diseases due to adults' vaccine refusal, or when older bikers inspire teenagers to risk their lives by riding without a helmet.

Driver Assistance and Self-driving Cars

Vehicle technology is in a process of increasing automation. Modern cars have a whole set of driver assistance functions that reduce the risk or the severity of accidents: anti-lock braking systems, cruise control, lane departure warning systems, driver attention monitors, collision warning and automatic brake systems, etc. This began with systems that merely warn the driver, but increasingly, functions are introduced that take control of the car if the driver does not heed a serious warning, for instance, automatic braking to avoid a head-on collision or mitigate its effects. There does not seem to have been much opposition to these systems. This may be because they have been constructed not to reduce the driver's feeling of control in normal driving.

However, it does not follow from these experiences that fully automatic, self-driving cars will easily be accepted. As it now seems, autonomous vehicles will not

be introduced, at least not on a large scale, until and unless they are far much safer than manually driven cars. We tend to be less willing to accept machine errors than errors made by human beings. Our tolerance for safety-critical malfunctions in cars is already low. Vehicles are recalled for repair even of faults that have a comparatively low probability of giving rise to injuries or fatalities. Occasionally, car manufacturers have refrained from recalls that they deemed too costly in relation to the gains in safety, but that led to strong negative reactions from the public (Smith 2017). Some authors have worried that a large number of avoidable fatalities and severe injuries will ensue if the introduction of autonomous vehicles is delayed by safety requirements that are inordinately high as compared to those applied to conventional vehicles (Brooks 2017; Hicks 2018, p. 67).

To the extent that automated vehicles run a smaller risk of crashes than humanly driven vehicles, there may be a movement towards relaxing safety measures such as the use of seat belts. Already in 2006, a year when more than 300 persons were killed in Swedish road traffic and around 9000 were injured, a Swedish physician (!) asked: “Do we even need the seat belts any more, now that the cars have become better?” (Eberhard 2006, p. 321) The answer to that question depends, of course, on your attitude to human death and suffering. Needless to say, less safe behavior by car occupants can at least partly offset the gains in safety obtainable with the new technology. Since no driver is needed, children can presumably travel alone in a self-driving car, and so can a company of severely inebriated persons. This can have consequences for occupant behavior. This means that questions about safety culture and how we influence each other’s behavior will continue to be pertinent in automated vehicles.

In a traffic system with automated vehicles that are much safer than human-driven cars, questions may arise about the legitimacy of the latter. Proposals have already been made to prohibit human driving on at least parts of the road net (Nyholm and Smids 2020). This can surely meet with ardent resistance, since driving is a major source of pleasure and pride for many people (Edensor 2004; Moor 2016; Borenstein et al. 2019, p. 392). However, such conflicts can possibly be avoided if future human-driven cars are equipped with advanced assistive technology, including an emergency autopilot function that takes over in extreme cases when this is necessary to avoid a crash.

Some authors have maintained that human moral agency is in some way stymied by technology that reduces the risk of injuries or other adverse effects of human failures. These authors call such technology “techno-regulation” (Brownsword 2005) or “technology paternalism” (Spiekermann and Pallas 2006). Examples that have been used include automated entrance barriers at railway and metro stations (to prevent fare dodging) (Yeung 2011, p. 22), alcohol interlocks (Spiekermann and Pallas 2006, p. 10), and – of course – self-driving cars (Brownsword 2005, pp. 16 and 18). Roger Brownsword maintains that these technologies have a problem in common:

When regulators design people, products, or places in such a way that regulatees have no choice but to act in whatever way is judged to be appropriate, we cannot meaningfully speak

about them being morally responsible agents, to be credited with acts that respect others and to be blamed where they fall short of moral requirements. (Brownsword 2005, p. 18)

Similarly, Karen Yeung writes:

Although an isolated techno-regulatory measure may appear benign, the *cumulative* effect of techno-regulatory action by a range of regulators acting independently across a variety of social contexts might ultimately lead to such a significant erosion of moral freedom that meaningful moral agency can no longer be sustained. (Yeung 2011, p. 27)

These arguments are based on an extremely one-sided selection of technological innovations. All through human history, technology and social arrangements have shifted in ways that have both created and closed down options for moral decision-making. New weapons have given rise to moral issues on when they could legitimately be used. Would it have been a disadvantage if weapons of mass destruction were not available, so that no one could exercise their moral agency by deciding (hopefully) not to use them? New means of food production have reduced the burden of deciding how to distribute food in times of famine. Would it have been better if there was even more scarcity of food so that more people could flourish as moral agents by distributing it rightfully? Technological devices, such as fences, locks, safes, and alarm systems have made burglary next to impossible in some places. Would it be preferable to ensure that all jewelry shops and bank vaults (and your home) have doors that are easy to force open, so that prospective burglars have ample opportunity to make virtuous decisions not to break in?

Arguably, these examples are sufficient to show the absurdity of valuing options for moral deliberation higher than the avoidance of human suffering. But it may nevertheless be of some interest to ask: Are these authors right in contending that modern technology, as a general tendency, removes moral choices and thereby moral agency from us humans? That is extremely difficult to determine, since examples pointing in both directions are easily found. Notably, the Internet as well as social media have given rise to a host of new moral issues that concern our private lives, information on climate change has put air travel in a new moral perspective, and environmental concerns have added moral aspects to our choices of foodstuff and other consumer products. This may very well add up to outweigh a potential reduction in the need for moral choices in road traffic.

In Conclusion: Vision Zero

Vision Zero has not been much mentioned in this chapter, but in a sense it has been all about Vision Zero. Questions about personal freedom and how we influence each other's behavior are central for the choice of a traffic safety strategy, and many of the differences between Vision Zero and previous road safety policies are closely connected with the issues we have penetrated. In at least three respects, the conclusions we have arrived at can serve as moral underpinnings of Vision Zero.

First, we found that some of the freedoms that were deeply cherished in former times, such as the freedoms to have “a little dirt” in your water pipes, to hire workers to work in mines with incessant rockfalls, or to drive drunk, are now considered by almost everyone as weird wishes, rather than worthy objects of liberty. This historical perspective shows that conceptions of what is normal and desirable can change. Vision Zero aims at one fundamental such change: Conditions and behaviors that lead to deaths and serious injuries in road traffic should no longer be seen as normal, or even acceptable.

Secondly, we noted that many habits and actions that are commonly conceived as “self-harming” are in fact the result of a combination of self-harming and other-harming actions. The smoker’s own decision is only part of the causal history behind her unhealthy habit. Another, in fact quite crucial, part is the tobacco industry’s ruthless promotion of a death-bringing product. Similarly, the motorist who speeds on the expressway is certainly to blame, but the reason why it is at all possible for him to do so is that the car manufacturer sells vehicles that can be driven at illegal speeds, and that the legislature allows this. This analysis shows that we need to consider the systemic causes behind events that have traditionally been ascribed exclusively to individual road users. This is one of the methodological cornerstones of Vision Zero.

Thirdly, we saw that traffic behavior is subject to strong effects of social contagion. For instance, people tend to wear seat belts and helmets if others do so, and to drive slower if others refrain from speeding. Therefore, safe traffic behavior is important not only for the direct effects of one’s own actions, but also as a contribution to the creation of a safety culture that protects us all. This is the major reason why driving a motorcycle without a helmet is not a “purely self-harming” action; it unavoidably sets a negative example that contributes to enticing other bikers to do the same. This is a communal and socially inclusive perspective on traffic safety that has not had a large role in Vision Zero, but it can arguably contribute constructively to its implementation.

Cross-References

► [Arguments Against Vision Zero: A Literature Review](#)

Acknowledgment I would like to thank Henok Girma Abebe, Karin Edvardsson Björnberg, Kalle Grill, Jesper Jerkert, Claes Tingvall and the participants in seminar at the Philosophy Division at KTH in October 2020 for valuable comments on an earlier version of this chapter.

References

- Anon. (1894). Chicago. *British Medical Journal*, 1(1726), 216–217.
- Apollonio, D. E., & Bero, L. A. (2007). The creation of industry front groups: The tobacco industry and ‘get government off our back’. *American Journal of Public Health*, 97(3), 419–427.

- Bachynski, K. E. (2012). Playing hockey, riding motorcycles, and the ethics of protection. *American Journal of Public Health*, 102(12), 2214–2220.
- Balko, R. (2011). Nanny state: Abolish drunk driving Laws. *Reason*, January 2011. Downloaded April 10, 2020, from <https://reason.com/2010/12/31/abolish-drunk-driving-laws>
- Barclift, P. L. (1991). In controversy with Saint Augustin: Julian of Eclanum on the nature of sin. *Recherches de Théologie Ancienne et Médiévale*, 58, 5–20.
- Barnard, J. (2002). I loved her because she wanted to climb the highest peak. That's who she was. *Guardian* 28 August.
- Borenstein, J., Hekert, J. R., & Miller, K. W. (2019). Self-driving cars and engineering ethics: The need for a system level analysis. *Science and Engineering Ethics*, 25(2), 383–398.
- Boughton, B. J. (1984). Compulsory health and safety in a free society. *Journal of Medical Ethics*, 10(4), 186–189.
- Brooks, R. (2017). Unexpected consequences of self driving cars. <https://rodneebrooks.com/unexpected-consequences-of-self-driving-cars>. Accessed 15 Apr 2020.
- Brownsword, R. (2005). Code, control, and choice: Why east is east and west is west. *Legal Studies*, 25(1), 1–21.
- Cairns, H. (1941). Head injuries in motorcyclists. The importance of the crash helmet. *British Medical Journal*, 2(4213), 465–471.
- Cairns, H. (1946). Crash helmets. *British Medical Journal*, 2(4470), 322–323.
- Camerer, C., Issacharoff, S., Loewenstein, G., O'donoghue, T., & Rabin, M. (2003). Regulation for conservatives: Behavioral economics and the case for 'asymmetric paternalism'. *University of Pennsylvania Law Review*, 151(3), 1211–1254.
- Cameron, M. H., Peter Vulcan, A., Finch, C. F., & Newstead, S. V. (1994). Mandatory bicycle helmet use following a decade of helmet promotion in Victoria, Australia – An evaluation. *Accident Analysis and Prevention*, 26(3), 325–337.
- Cardador, M. T., Hazan, A. R., & Glantz, S. A. (1995). Tobacco industry smokers' rights publications: A content analysis. *American Journal of Public Health*, 85(9), 1212–1217.
- Carson, R. (1962). *Silent spring*. Boston: Houghton Mifflin.
- Carsten, O. M. J., & Tateb, F. N. (2005). Intelligent speed adaptation: Accident savings and cost-benefit analysis. *Accident Analysis and Prevention*, 37, 407–416.
- Chancellor, A. (2003). In need of speed. *The Guardian*, 7 June.
- Christakis, N. A., & Fowler, J. H. (2013). Social contagion theory: Examining dynamic social networks and human behavior. *Statistics in Medicine*, 32(4), 556–577.
- Chrysostomus, J. (1836). Diatriba ad opus imperfectum in Matthaeum. In *Opera Omnia*, Tomus Sextus (pp. 731–972). Paris: Gaume Frates.
- Conly, S. (2017). Paternalism, coercion and the unimportance of (some) liberties. *Behavioural Public Policy*, 1(2), 207–218.
- Connolly, T., & Åberg, L. (1993). Some contagion models of speeding. *Accident Analysis & Prevention*, 25(1), 57–66.
- Crossley, D. (1999). Paternalism and corporate responsibility. *Journal of Business Ethics*, 21(4), 291–302.
- Daker, G. et al. (1893). *Foreign Abstract of Sanitary Reports*, 8(30), 630–648.
- de Bellis, E., Schulte-Mecklenbeck, M., Brucks, W., Herrmann, A., & Hertwig, R. (2018). Blind haste: As light decreases, speeding increases. *PLoS One*, 13(1), e0188951.
- de Marneffe, P. (2006). Avoiding paternalism. *Philosophy and Public Affairs*, 34(1), 68–94.
- Delaney, A., Ward, H., Cameron, M., & Williams, A. F. (2005a). Controversies and speed cameras: Lessons learnt internationally. *Journal of Public Health Policy*, 26(4), 404–415.
- Delaney, A., Ward, H., & Cameron, M. (2005b). *The history and development of speed camera use* (Report no. 242). Monash University Accident Research Centre.
- Dorrian, C. I. (1911). Better automobile laws. *Country Life in America*, 19(12), 467.
- Dworkin, G. (1972). Paternalism. *The Monist*, 56(1), 64–84.
- Dworkin, G. (2015). Defining paternalism. In T. Schramme (Ed.), *New perspectives on paternalism and health care* (pp. 17–29). Cham: Springer.

- Dworkin, G. (2020). Paternalism. In E. N. Zalta (Ed.), *The Stanford encyclopedia of philosophy*. <https://plato.stanford.edu/archives/spr2020/entries/paternalism/>.
- Eberhard, D. (2006). *I trygghetsnarkomanernas land: om Sverige och det nationella paniksyndromet*. [In the land of security addicts. On Sweden and the national panic syndrome.] Stockholm: Prisma.
- Edensor, T. (2004). Automobility and national identity: Representation, geography and driving practice. *Theory, Culture and Society*, 21(4–5), 101–120.
- Edwards, J., Freeman, J., Soole, D., & Watson, B. (2014). A framework for conceptualising traffic safety culture. *Transportation Research Part F*, 26, 293–302.
- Elston, P. A. (1971). *Maximum speed limits. The development of speed limits: a review of the literature*. Report under contract FH-11-7275 for the National Highway Traffic Safety Administration. Indiana University, Institute for Research in Public Safety. Downloaded April 10, 2020 from <https://rosap.nhtl.bts.gov/view/dot/16506>
- Elvebakk, B. (2015). Paternalism and acceptability in road safety work. *Safety Science*, 79, 298–304.
- Evans, L. (1991). *Traffic safety and the driver*. New York: Van Nostrand Reinhold.
- Fallin, A., Grana, R., & Glantz, S. A. (2014). ‘To quarterback behind the scenes, third-party efforts’: The tobacco industry and the tea party. *Tobacco Control*, 23(4), 322–331.
- Farmer, C. M. (2017). Relationship of traffic fatality rates to maximum state speed limits. *Traffic Injury Prevention*, 18(4), 375–380.
- Feinberg, J. (1986). *Harm to self*. Oxford: Oxford University Press.
- Fell, J. C., & Voas, R. B. (2014). The effectiveness of a 0.05 blood alcohol concentration (BAC) limit for driving in the United States. *Addiction*, 109(6), 869–874.
- Fell, J. C., Lacey, J. H., & Voas, R. B. (2004). Sobriety checkpoints: Evidence of effectiveness is strong, but use is limited. *Traffic Injury Prevention*, 5(3), 220–227.
- Flanigan, J. (2017). Seat belt mandates and paternalism. *Journal of Moral Philosophy*, 14(3), 291–314.
- Fu, C., Liu, H., & Zhou, Y. (2020). A review of the speeding intervention effectiveness and acceptance of intelligent speed adaptation. In *Sixth International Conference on Transportation Engineering, Proceedings (ICTE 2019)* (pp. 299–310). Reston: American Society of Civil Engineers.
- Gabriel, R. A. (2007). *Soldiers’ lives through history: The ancient world*. Westport: Greenwood Press.
- Gardner, E. (1941). Head injuries in motor-cyclists. *British Medical Journal*, 2(4216), 592–592.
- Giubilini, A., & Savulescu, J. (2019). Vaccination, risks, and freedom: The seat belt analogy. *Public Health Ethics*, 12(3), 237–249.
- Green, T. H. (1881). Lecture on liberal legislation and freedom of contract. In R. L. Nettleship (Ed.), *Works of Thomas Hill Green* (Vol. 3, pp. 365–386). London: Longmans, Green and Co.
- Green, E. H. (1995). *The crisis of conservatism. The politics, economics and ideology of the British conservative party, 1880–1914*. London: Routledge.
- Grill, K. (2018). Paternalism by and towards groups. In K. Grill & J. Hanna (Eds.), *The Routledge handbook of the philosophy of paternalism* (pp. 46–58). Routledge: Abingdon.
- Grill, K., & Nihlén Fahlquist, J. (2012). Responsibility, paternalism and alcohol interlocks. *Public Health Ethics*, 5(2), 116–127.
- Hansson, S. O. (1999). The moral significance of undetectable effects. *Risk*, 10, 101–108.
- Hansson, S. O. (2005). Extended antipaternalism. *Journal of Medical Ethics*, 31, 97–100.
- Hansson, S. O. (2014). Making road traffic safer: Reply to Ori. *Philosophical Papers*, 43, 365–375.
- Hicks, D. J. (2018). The safety of autonomous vehicles: Lessons from philosophy of science. *IEEE Technology and Society Magazine*, 37(1), 62–69.
- House of Commons. (1847). The health of towns bill. Debate 18 June. *Hansard’s Parliamentary Debates*, 3rd series, Vol. 93, cc. 727–753.
- House of Commons. (1848). Public Health Bill. Debate 5 May. *Hansard’s Parliamentary Debates*, 3rd series, Vol. 98, cc. 710–743 and 764–803.

- Howat, P., O'Connor, J., & Slinger, S. (1992). Community action groups and health policy. *Health Promotion Journal of Australia*, 2(3), 16–22.
- Huguelet, A. (2020). 'It's about freedom': Local lawmaker trying again on motorcycle helmet law repeal. *Springfield News-Leader*, February 29, 2020. Downloaded April 5, 2020 from <https://eu.news-leader.com/story/news/politics/2020/02/29/motorcycle-helmet-law-repeal-rides-again-its-freedom/4902534002/>
- Hundley, J., et al. (2004). Non-helmeted motorcyclists: A burden to society? A study using the National Trauma Data Bank. *Journal of Trauma: Injury, Infection, and Critical Care*, 57(5), 944–949.
- Husak, D. N. (1994). Is drunk driving a serious offense? *Philosophy and Public Affairs*, 23(1), 52–73.
- Ichikawa, M., Nakahara, S., & Wakai, S. (2002). Mortality of front-seat occupants attributable to unbelted rear-seat passengers in car crashes. *Lancet*, 359(9300), 43–44.
- Jamrozik, E., Handfield, T., & Selgelid, M. J. (2016). Victims, vectors and villains: Are those who opt out of vaccination morally responsible for the deaths of others? *Journal of Medical Ethics*, 42(12), 762–768.
- Johannessen, G. (1984). *Historical perspective on seat belt restraint systems* (SAE Technical Paper no. 840392). Society of Automotive Engineers.
- Jones, M. M., & Bayer, R. (2007). Paternalism & its discontents: Motorcycle helmet laws, libertarian values, and public health. *American Journal of Public Health*, 97(2), 208–217.
- Katz, J. E. (2005). Individual rights advocacy in tobacco control policies: An assessment and recommendation. *Tobacco Control*, 14(Suppl. 2), ii31–ii37.
- Knowles, T. (1880). Speech in the House of Commons. *Hansard's Parliamentary Debates*, 3rd series, Vol. 252, cc. 1096–1102.
- Krupke, C. H., Prasad, R. P., & Anelli, C. M. (2007). Professional entomology and the 44 noisy years since Silent Spring. Part 2: Response to Silent Spring. *American Entomologist*, 53(1), 16–26.
- Langland-Orban, B., & Flint, L. (2011). Another perspective on motorcycle helmet use. *Journal of the American College of Surgeons*, 213(2), 336–337.
- Lanska, D. J. (2009). Historical perspective: Neurological advances from studies of war injuries and illnesses. *Annals of Neurology*, 66(4), 444–459.
- Le Grand, J., & New, B. (2015). *Government paternalism: Nanny state or helpful friend?* Princeton: Princeton University Press.
- Lehmann, S., & Ahn, Y.-Y. (Eds.). (2018). *Complex spreading phenomena in social systems influence and contagion in real-world social networks*. Cham: Springer.
- Leichter, H. M. (1991). *Free to be foolish*. Princeton: Princeton University Press.
- Liang, H., & Zhang, S. (2019). Impact of supervisors' safety violations on an individual worker within a construction crew. *Safety Science*, 120, 679–691.
- Liang, H., Lin, K.-Y., Zhang, S., & Yikun, S. (2018). The impact of coworkers' safety violations on an individual worker: A social contagion effect within the construction crew. *International Journal of Environmental Research and Public Health*, 15(4), 773.
- Marczinski, C. A., & Fillmore, M. T. (2009). Acute alcohol tolerance on subjective intoxication and simulated driving performance in binge drinkers. *Psychology of Addictive Behaviors*, 23(2), 238–247.
- Max, J. (2019). Speed cameras vulnerable to vandalism worldwide – But studies say they save lives. *Forbes*, 21 January. Downloaded 4 April 2020 from <https://www.forbes.com/sites/joshmax/2019/01/21/speed-cameras-vulnerable-to-vandalism-worldwide-but-studies-say-they-save-lives/#40f342b66f00>
- Mays, M. (2019). Anti-vaccine protesters are likening themselves to civil rights activists. *Politico*, 19 September. Downloaded April 5, 2020 from <https://www.politico.com/story/2019/09/18/california-anti-vaccine-civil-rights-1500976>
- McGhie, A., Lewis, I., & Hyde, M. K. (2012). The influence of conformity and group identity on drink walking intentions: Comparing intentions to drink walk across risky pedestrian crossing scenarios. *Accident Analysis and Prevention*, 45, 639–645.

- McKenna, F. P. (2007). The perceived legitimacy of intervention: A key feature for road safety. In *Improving traffic safety culture in the United States: The journey forward* (pp. 165–175). Washington, DC: AAA Foundation for Traffic Safety.
- Meyer, C., & Reiter, S. (2004). Impfgegner und Impfskeptiker. Geschichte, Hintergründe, Thesen, Umgang. *Bundesgesundheitsblatt – Gesundheitsforschung – Gesundheitsschutz*, 47, 1182–1188.
- Milio, N. (1976). A framework for prevention: Changing health-damaging to health-generating life patterns. *American Journal of Public Health*, 66, 435–439.
- Mill, J. S. (1977 [1859]). *On liberty*. In J. M. Robson (Ed.), *Collected works of John Stuart Mill* (Vol. 18, pp. 213–310). Toronto: University of Toronto Press.
- Miller, J. M. (1993). A swinging pendulum or Damocles' sword: A dramatic perspective on constitutional liberties in the 1990s. *Western State University Law Review*, 21(1), 165–218.
- Mols, F., Alexander Haslam, S., Jetten, J., & Steffens, N. K. (2015). Why a nudge is not enough: A social identity critique of governance by stealth. *European Journal of Political Research*, 54(1), 81–98.
- Moor, R. (2016). What happens to American myth when you take the driver out of it? The self-driving car and the future of the self. *Intelligencer, New York Magazine*, October 16. <http://nymag.com/selectall/2016/10/is-the-self-driving-car-un-american.html>
- Mowat, L. (2019). Nanny state: Speed limiter tech for your car has black box recording your EVERY move. *Daily Mail*, 28 March.
- National Motorists Association. (2019). Vision zero invasion of the car itself. *NMA E-Newsletter* #532. Downloaded 10 April 2020 from <https://www.motorists.org/blog/vision-zero-invasion-of-the-car-itself-nma-e-newsletter-532/>
- NHTSA. (2019). 2017 *State traffic data*. Traffic Safety Facts, DOT HS 812 780. Downloaded 19 April 2020 from <https://crashstats.nhtsa.dot.gov/#/>
- Nolte, K. B., Healy, C., Rees, C. M., & Sklar, D. (2017). Motorcycle policy and the public interest: A recommendation for a new type of partial motorcycle helmet law. *Journal of Law, Medicine and Ethics*, 45(1 Suppl), 50–54.
- Nyholm, S., & Smids, J. (2020). Automated cars meet human drivers: Responsible human-robot coordination and the ethics of mixed traffic. *Ethics and Information Technology*, 22(4), 335–344.
- O'Madagain, C. (2014). Can groups have concepts? Semantics for collective intentions. *Philosophical Issues*, 24(1), 347–363.
- Petroski, H. (2016). The lines on the road: Infrastructure in perspective. *Mechanical Engineering*, 138(2), 42–47.
- Pierik, R. (2018). Mandatory vaccination: An unqualified defence. *Journal of Applied Philosophy*, 35(2), 381–398.
- Pilkington, P. (2003). Speed cameras under attack in the United Kingdom. *Injury Prevention*, 9, 293–294.
- Porter, D. (1999). *Health, civilisation and the state. A history of public health from ancient to modern times*. London: Routledge.
- Preyer, G., Hindriks, F., & Chant, S. R. (Eds.). (2014). *From individual to collective intentionality: New essays*. Oxford: Oxford University Press.
- Rigau-Pérez, J. G. (1989). The introduction of smallpox vaccine in 1803 and the adoption of immunization as a government function in Puerto Rico. *Hispanic American Historical Review*, 69(3), 393–423.
- Roberts, D. (1958). Tory paternalism and social reform in early Victorian England. *American Historical Review*, 63(2), 323–337.
- Roberts, D. (1979). *Paternalism in early Victorian England*. New Brunswick: Rutgers University Press.
- Robinson, M. (2019). Gilets jaunes protesters vandalize 60% of France's speed cameras. *CNN*, January 11. <https://edition.cnn.com/2019/01/10/europe/gilets-jaunes-speed-cameras-destroyed-france-scli-intl/index.html>. Accessed 17 Oct 2020.

- Rosenberg, B., & Levenstein, C. (2010). Social factors in occupational health: A history of hard hats. *New Solutions: A Journal of Environmental and Occupational Health Policy*, 20(2), 239–249.
- Salice, A. (2015). There are no primitive we-intentions. *Review of Philosophy and Psychology*, 6(4), 695–715.
- Schneider, N. K., & Glantz, S. A. (2008). ‘Nicotine Nazis strike again’: A brief analysis of the use of Nazi rhetoric in attacking tobacco control advocacy. *Tobacco Control*, 17(5), 291–296.
- Schonsheck, J. (1994). *On criminalization. An essay in the philosophy of the criminal law*. Dordrecht: Kluwer Academic Publishers.
- Schramme, T. (2015). Contested services, indirect paternalism and autonomy as real liberty. In T. Schramme (Ed.), *New perspectives on paternalism and health care* (pp. 101–114). Cham: Springer.
- Sheikh, A., Cook, A., & Ashcroft, R. (2004). Making cycle helmets compulsory: Ethical arguments for legislation. *Journal of the Royal Society of Medicine*, 97(6), 262–265.
- Shiffrin, S. V. (2000). Paternalism, unconscionability doctrine, and accommodation. *Philosophy and Public Affairs*, 29(3), 205–250.
- Smids, J. (2018). The moral case for intelligent speed adaptation. *Journal of Applied Philosophy*, 35(2), 205–221.
- Smith, F. F. (1967). *Controlling insects on flowers* (Agricultural Information Bulletin, no. 237). U.S. Department of Agriculture.
- Smith, B. W. (2017). The trolley and the Pinto: Cost-benefit analysis in automated driving and other cyber-physical systems. *Texas A&M Law Review*, 4, 197–208.
- Snell, L. M. (2018). Hard hat – Origins and evolution. *Concrete International*, 40(7), 48–49.
- Solan, D. (1986). Seat-belt laws violate your civil rights. *New York Times*, 26 February. National edition, Section A, p. 22.
- Solomon, D. (2010). Questions for Rand Paul. Tea time. *New York Times Magazine*, 31 March.
- Spiekermann, S., & Pallas, F. (2006). Technology paternalism—wider implications of ubiquitous computing. *Poiesis and Praxis*, 4(1), 6–18.
- Spurgin, E. W. (2006). Occupational safety and paternalism: Machan revisited. *Journal of Business Ethics*, 63(2), 155–173.
- Stonehouse, A. (2017). Speech in the Legislative Council of the Parliament of Western Australia, 24 May 2017. Downloaded 5 April 2020, from [https://www.parliament.wa.gov.au/parliament/Memblast.nsf/\(MemberPics\)/8D0A3D52071587BF482580F30020A00C/\\$file/AaronStonehouseInaug2017.pdf](https://www.parliament.wa.gov.au/parliament/Memblast.nsf/(MemberPics)/8D0A3D52071587BF482580F30020A00C/$file/AaronStonehouseInaug2017.pdf)
- Thaler, R. H., & Sunstein, C. R. (2008). *Nudge: Improving decisions about health, wealth, and happiness*. New Haven: Yale University Press.
- Tripp, H. A. (1928). Police and public: A new test of police quality. *Police Journal*, 1(4), 529–539.
- Underwood, H. (1907). Speed mania and how to cure it. *Harper Weekly*, 51:1(2623), 470.
- Vallgård, S. (2012). Nudge – A new and better way to improve health? *Health Policy*, 104(2), 200–203.
- Vissing, J. (2014). Preserving liberty while increasing safety: Why Indiana should outlaw sobriety checkpoints. *Indiana Law Review*, 48, 1089–1113.
- Vivoda, J. M., & Eby, D. W. (2011). Factors influencing safety belt use. In B. E. Porter (Ed.), *Handbook of traffic psychology* (pp. 215–230). Waltham: Academic Press.
- Waller, G. (1988). Multiple contributions to a debate on the Road Traffic Act, *Hansard's Parliamentary Debates*, series 6, vol. 133, cc. 588–621, 13 May.
- Waller, P. F. (2001). Public health's contribution to motor vehicle injury prevention. *American Journal of Preventive Medicine*, 21(4), 3–4.
- White, D. (2006). Take two servings of paternalism. Downloaded April 5, 2020, from <https://www.aei.org/articles/take-two-servings-of-paternalism/>
- Wiggins, D. E. (1987). *The cholera and public health in Mid-Victorian Britain*. Master's thesis, Texas Tech University. Downloaded 25 April 2020 from <https://ttu-ir.tdl.org/bitstream/handle/2346/14620/31295005215370.pdf?sequence=1>

- Williams, A. F. (2006). Alcohol-impaired driving and its consequences in the United States: The past 25 years. *Journal of Safety Research*, 37(2), 123–138.
- Wilson, J. (2015). Why it's time to stop worrying about paternalism in health policy. In T. Schramme (Ed.), *New perspectives on paternalism and health care* (pp. 203–218). Cham: Springer.
- Woolfson, D. (2016). Barbecue pub slams 'nanny state' after zero food hygiene rating, 29 August 2016. Downloaded April 5, 2020, from <https://www.morningadvertiser.co.uk/Article/2016/08/30/BBQ-pub-slams-nanny-state-after-zero-food-hygiene-rating>
- Yerger, V. B., Przewoznik, J., & Malone, R. E. (2007). Racialized geography, corporate activity, and health disparities: Tobacco industry targeting of inner cities. *Journal of Health Care for the Poor and Underserved*, 18(6), 10–38.
- Yeung, K. (2011). Can we employ design-based regulation while avoiding Brave New World? *Law, Innovation and Technology*, 3(1), 1–29.
- Zaidel, D. M. (1992). A modeling perspective on the culture of driving. *Accident Analysis and Prevention*, 24(6), 585–597.

Open Access This chapter is licensed under the terms of the Creative Commons Attribution 4.0 International License (<http://creativecommons.org/licenses/by/4.0/>), which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if changes were made.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

