

Intertheoretic Relations in Context: Details, Purpose, and Practice

Joseph T. Ricci

A dissertation submitted in partial fulfillment of the requirements for the degree of

Doctor of Philosophy

University of Washington
2015

Reading Committee:
Andrea Woody, Chair
Arthur Fine
John Manchak
Alison Wylie

Program Authorized to Offer Degree:
Philosophy

© Copyright 2015
Joseph T. Ricci

Abstract

Intertheoretic Relations in Context: Details, Purpose, and Practice

Joseph T. Ricci

Chair of the Supervisory Committee:

Andrea Woody, Associate Professor

Philosophy

An intertheory comparison should be assessed with regards to what goals it seeks to accomplish. Traditionally reductions have sought to establish ontological primacy, and also to have the reducing theory explain features of the reduced. From a functionalist perspective, this dissertation assesses three major results: a reduction of general relativistic spacetimes to a Newtonian gravitational structure, a theory comparison that employs both wave and ray optics, and a reduction that limits the momentum equation of special relativity to the classical momentum equation. These case-studies reveal evidence of types of goals given little discussion in existing literature on theory reduction. I find that successional reductions can (i) provide an explanation of (aspects of) the succeeded theory by the successor; (ii) provide an explanation of the theories' successes/failures, as well as explaining details of the progress, both historical and conceptual, from the succeeded theory to the successor; (iii) transfer confidence to the successor from the succeeded theory; and (iv) delimit a range of applicability for the succeeded theory. Recognizing these new goals provides insight for analysis of the intertheoretic activity of scientists, as well as the work of philosophers that assess how scientific theories relate to one another.

Table of Contents

Acknowledgements.....	i
Dedication.....	iii
General Introduction.....	1
 Chapter 1 - Capturing Reduction by Models: Early Attempts to Compare Theories	
1.0 Introduction.....	5
1.1 Three Historical Models of Reduction.....	6
1.1.1 Nagel's Deductive Model.....	7
1.1.2 Kemeny and Oppenheim's Disjoint-Explanation Model.....	15
1.1.3 Suppes's Semantic-Isomorphism Model.....	26
1.2 Goals and Aims of Philosophical Models of Reduction.....	32
1.2.1 Two Traditional Goals of Reduction.....	32
1.2.2 Relata and Intention.....	34
1.3 Conclusion.....	36
 Chapter 2 - Analogous Theories: General Relativity and Newtonian Mechanics	
2.0 Introduction.....	37
2.1 The Schaffner Model.....	38
2.1.1 The New Wave Model.....	42
2.2 Trautman's Reduction.....	48
2.2.1 Traditional Epistemic Goals: Theory Explaining Theory.....	55
2.2.2 Differing Epistemic Goals: Reduction Explaining Scientific Progress.....	57
2.3 Conclusion.....	60
 Chapter 3 - Limits and Approximations: Special Relativity and Classical Mechanics	
3.0 Introduction.....	62
3.1 Nickles's Two Models for Reduction.....	63
3.2 The Problem of Limiting c	66
3.2.1 Explaining Scientific Progress.....	72
3.3 The Problem of Limiting v	73
3.3.1 Success and Context.....	76
3.3.2 Transference of Confidence.....	81
3.4 The Problem of Limiting $(v/c)^2$	83
3.4.1 Establishing Past Theories and Applicability.....	89
3.4.2 Recasting Prior Successes.....	92
3.5 Conclusion.....	95

Chapter 4 - Intra-level Relations: Wave Optics and Ray Optics

4.0 Introduction.....	96
4.1 Batterman and Limiting.....	97
4.2 Catastrophe Optics.....	99
4.3 Belot's <i>Ab Initio</i> Objection.....	107
4.4 Batterman's Contextual Response.....	111
4.5 Redhead's Accusation of Reification.....	113
4.5.1 Intertheory Comparisons as Scientific Activity.....	118
4.6 Conclusion.....	120

Chapter 5 - Reductive Models, Ordering, and Scientific Structure

5.0 Introduction.....	123
5.1 Models of Reduction and Three Ordering Relations.....	124
5.1.1 Nagel's Deductive Model.....	126
5.1.2 Kemeny and Oppenheim's Disjoint-Explanation Model.....	133
5.1.3 Suppes's Semantic-Isomorphism Model.....	136
5.1.4 Schaffner's Model.....	138
5.1.5 The New Wave Model.....	142
5.2 Orderings of Science.....	147
5.2.1 Reflexivity and Triviality.....	148
5.2.2 Symmetry and Equivalence.....	151
5.2.3 Transitivity and Succession.....	155
5.2.4 Ranking Theories to Make a Structure for Science.....	157
5.3 Conclusion.....	161

Chapter 6 - Functionalism and Goals

6.0 Introduction.....	163
6.1 Reasons for Comparing Theories.....	163
6.1.1 Pluralism and Fundamentalism.....	168
6.1.2 Aristotelian Dynamics and Newtonian Mechanics.....	169
6.2 Scerri's Assessment of Quantum Chemistry.....	173
6.2.1 The Functionalist Therapy.....	178
6.3 Contributions and Conclusions.....	181

Appendix

G Glossary.....	183
-----------------	-----

Bibliography

B Bibliography.....	185
---------------------	-----

Acknowledgements

Let me begin by introducing the real hero of this story: Andrea Woody. In 2004, I met in her office for the first time to discuss my interest in going to graduate school. She later wrote me to tell me I was accepted to UW. Over the 10+ years since then, she has been a constant source of support, knowledge, and inspiration. She is the chair of my committee, my teacher, my confidant, and the wisest person I know – thank you for believing I could do this Andrea.

The Philosophy Department at UW has been a central fixture of my life here in Seattle. The many professors and staff members have always been encouraging and important to making me feel like I belong. My committee members Arthur Fine, Alison Wyle, and Jon Manchak were all essential to making this document happen. I would particularly like to thank Ann Baker and Barbara Mack for always being there when I needed advice. Many of the graduate students, past and present, have been indispensable in giving guidance and tolerating my naïve questions. I would like to especially notice Elizabeth Scarbrough, Rachel Fredericks, and Jon Rosenberg. Most important in this group is Brandon Morgan-Olsen, supporting me as an academic, as a colleague, and as a friend. I would also like to thank my students – over 1,000 of them – for believing that I was worth listening to. Through all of the people involved, I look back at my time teaching, studying, and working at UW with much happiness.

I wear many hats, and as such there are many people who have been there to help make this dissertation possible. To my Seattle network of friends, thank you for many crazy times. Sean Bray, Tim Uomoto, and Jon Francois all have especially accepted me at my best and at my worst. I would like to thank my co-workers and students at Seattle Central College for reassuring me that I make a difference. To Kirby Green and the rest of the Miss Sherri crew – thanks for not giving me weird looks when I had to work on this dissertation during my wheel watches. I'd like to thank Matt Baker working around my academic schedule to travel, and for waking up early to practice before my classes. Thanks to all the bboys in the Seattle dancing community, especially Circle of Fire, for reminding me that in addition to being an academic I am also an artist. Were it not for all of the great people I have known in Seattle, I would not have been able to preserve with this process.

Finally I would like to thank my family. My mother and father have always supported my academic interests without pressure and without judgment. They cultivated my interests,

allowing me to argue over dinner and read late in to the night with books they bought for me. Knowing that I was accepted regardless of my choice of major, my career path, or whether I decided to finish my PhD allowed me to work without any of the worries that many others are plagued by. It was my father who let me take home *The Portable Nietzsche* when I was 14, a book that began my philosophical journey. My brother, who has unfortunately gotten his PhD before me, has always been there when I needed someone the most, whether being fast conversation over the phone or taking me in over the weekend when I appeared after things fell apart. James, Mom, and Dad – I am finally finished.

Dedication

To the UW Philosophy Department

General Introduction

Scientists have long compared theories to one another, and with many purposes in mind. This dissertation brings attention to these multifarious goals, demonstrating their important role in a philosophical analysis of the comparison of theories that border one another, be it temporally, mereologically, or thematically. By being aware of the historical and ideological progression of the philosophical literature on the subject, this dissertation will follow philosophers of science as they focus on: theory reductions, as they relate to the unity of science; successional reductions, attending to the trajectory of science; and finally, asymptotic analyses, inasmuch as they facilitate an understanding of a theory's scope and domain of applicability. A major theme will be to observe how different theory comparisons appear to satisfy quite distinct aims, and consequently, that philosophical models must proceed very cautiously when attempting to characterize intertheoretic behavior generally. "Reductions", "relations", and their ilk have different meanings in different contexts. This being the case, recognition of the various motivations resting behind these theory comparisons, some of which have gone otherwise overlooked, will ultimately provide a commonality amongst the differing modes of inter-theory activity: each philosophical analysis should keep the goals of the endeavor in mind. Although less ambitious than an all-encompassing model for theory relations, this thesis facilitates an understanding of the general motivations for scientists reducing theories, relating theories, or asymptotically joining theories, while also revealing useful commonalities.

The first chapter focuses on one of the earliest and most discussed modes of intertheoretic activity: reduction of one theory to another. In the chapter, I discuss the differences and merits of three conceptually diverse models of theory reduction that are also historically significant, calling attention to features and themes that will be relevant throughout this dissertation. One major theme of the chapter concerns *scope*: attention to the type of theories that may be considered by a model of reduction. The earlier archetypes attempted to be applicable to all theories, reducing macroscopic to microscopic, old to new, and less-fundamental to more-fundamental. In addition, the chapter will show how each of the three early models of theory reduction demanded an exact correlation between particular claims of the reducing theory and

those of the reduced. The reducing theory may, however, expand to include claims outside the purview of the reduced theory.

The second chapter begins by examining how more articulated models of reduction accommodate for scientific progress in the light of failure. Often we find that theories improve upon the predictions of other theories, and the two are only *approximately* in agreement with one another. The principle maneuver is to employ the reducing theory to create an *analog* theory, one that attempts to mimic the predictions and mechanisms of the reduced theory proper. In this way, the falsity of the reduced theory is given provision by means of this analogy. As analogies may be strong or weak, so may reductions of this sort. In this framework, it becomes manifest that the degree of an analogy warrants different inferences to be drawn in different cases, for example, retention of ontology in some cases, or replacement of theoretical mechanisms in others.

The final part of the second chapter provides a case study of the reduction of general relativity to Newtonian mechanics. As a test case, this example provides excellent fodder to attempt a fit to the models of approximate reduction examined earlier in the chapter. After an examination of the scientific details of the reduction, in relation to the philosophical reduction models, the discussion turns to focus on what has been accomplished by the reduction. The goals of this successional reduction are twofold: first, that general relativity can explain features of Newtonian mechanics that were otherwise puzzling and in need of explanation; secondly, that the reduction provides an explanation of why Newton's theory was considered the dominant physical description of our world for so long – of why it was so successful.

The third chapter examines philosophers' attempts to characterize the role of limiting in successional reductions, from the perspective of a single case study. The limiting of the momentum equation of special relativity to the momentum equation of classical mechanics has received significant attention in the literature, and for some is seen as a paradigmatic example of the role limiting plays in reductions. By looking closely at the details of the case, it becomes clear that a mathematical limiting relationship on its own is impoverished. First, there are cases of limiting that seem poor candidates for reductions, on the basis of how poorly the equations represent empirical phenomenon. Secondly there are other examples where a reduction seems quite justified, in regards to how well they accomplish the goals of successional reductions, yet fail to exhibit a limiting relationship. One of the resounding lessons of the chapter is the important role played by the physical circumstances of values that are being limited. What the

values represent empirically, as well as what the limited quantity will signify in the world, are essential details that cannot be ignored. Finally, the specific benefits of the case study emerge in the form of new goals: a successional reduction may explain the successes of a prior theory, it may transfer confidence into a new theory, and lastly it may help to clarify the role played by past theories in the progression of science.

The fourth chapter examines an intra-level case study that aims to describe certain universal features of the rainbow. It begins by examining the claim that a proper explanation of certain rainbow phenomena requires resources from both wave optics, the currently-accepted theory, and ray optics, the theory it succeeded. The details of such an account at times treat light as a wave, so as to account for interference, and at other times treat light as a ray, so as to account for variations in properties such as the size of raindroplets. The chapter proceeds to consider an objection to the indispensability of ray optics in the scientific treatment of the phenomena. A pure mathematician, operating solely from wave-theoretic equations, could arguably achieve the desired results. The issue hinges on whether the mathematician tacitly relies on ray theoretic information to impose constraints on the wave-theoretic equations. The function of ray theory in the discussion, *qua* actualized scientific theory or *qua* unrealized mathematical aide, is ultimately left an open question. Moreover, while this extended scientific example involves asymptotic analysis between succeeding and succeeded theories, the goals of the case are merely to provide an adequate theoretical understanding of empirical phenomena. Thus the case-study should not be viewed as a reduction, but instead as an intra-level comparison.

In the fifth chapter, we return to examine certain logical features of the five models of reduction introduced in the first two chapters. This analysis concerns the possible ways that we could construe the order of the constituent theories involved in a reduction. Specifically, we will find for each reduction model the possibilities of reflexivity, antisymmetry, and transitivity. Interestingly, some models are reflexive while others are not, some models allow for symmetric cases, and not all models are transitive in all cases. Next, the three logical relations are discussed in light of possible goals for a reduction. Reflexivity is found to be a relation of little value, as its existence/non-existence has no bearing on any of the possible goals for a reduction. Antisymmetry is important for charting intra-level progress over time, as well as for tracking inter-level compositionality. Symmetric cases, on the other hand, provide examples where two theories might be used interchangeably. Transitivity provides the ability to relate theories that are

far away from one another along inter-level or intra-level chains. This chapter concludes by considering how antisymmetric and transitive reductions may provide an architectonic for science. In some cases, an ordering of theories may generate the fundamentalist's scientific pyramid, but I contend it may also result in a partially-ordered patchwork of theories more amenable to the pluralist's worldview.

The final chapter draws general conclusions based on the work of the previous five chapters. It summarizes the goals that have been associated with intertheory comparisons, and underscores that the view on reduction offered in this dissertation is thoroughly functionalist. Considering the work done discussing reductions and ordering of science in the fifth chapter, the sixth chapter purports that earlier results from chapter four can be seen as an argument in favor of pluralism. The chapter also suggests the value of a genuine transitive reduction – Aristotelian dynamics to Newtonian mechanics to general relativity – claiming that the reduction would provide benefits beyond those associated with the case as a mere historical curiosity. This argument can be made succinctly once we consider the scientific examples from chapters two and three, as well as the work on transitivity done in chapter five.

Next, the chapter provides a brief example of how a functional sensibility may elucidate existing discussions of reduction, looking at a philosophical paper that examines a case from quantum chemistry. Once it is admitted that reduction may accomplish many goals, the value of the scientific work originally dismissed here is showcased in a new light. Finally, the chapter closes with a summary of the novel contributions that this dissertation has made to the existing literature.

Chapter 1 - Capturing Reduction by Models:

Early Attempts to Compare Theories

Questions about reduction – what is its nature, and whether it is possible at all – are much more subtle than they are often taken to be.

Robert Batterman (2002a, 6)

§1.0 Introduction:

This chapter focuses on reduction, detailing how it has been considered by philosophers in the past: what models they have put forth and how these models relate to goals one might have for a reduction. §1.1 introduces three major conceptual models of reduction that are also of historical importance. §1.1.1 examines a deductive account conceived by Ernest Nagel, while §1.1.2 details an explanatory model developed by J. G. Kemeny and P. Oppenheim. §1.1.3 presents a model attributed to a remark made by Patrick Suppes, showing how we could view reduction as an isomorphism between models. In each subsection, I give an analysis of the strengths and weaknesses inherent to each model. Also, the sections will feature a discussion on what each model allows for the *relata* of their reduction relation: whether theories, theory parts, models, etc. are allowed for consideration in a reduction.

§1.2 draws conclusions from the attention paid to the *relata*. To reach these conclusions we must first, in §1.2.1, examine the goals one might have for pursuing a reduction. This brief discussion will allow us to distinguish between *ontological goals* and *epistemic goals*. §1.2.2 will conjecture about what the allowed *relata* reveal about the intentions of the three authors considered, as well as examining the scope of each model. This discussion will also allow us to begin to see the shortcomings of the three models, setting the stage for a transition into the three case studies that occupy Chapters §2, §3, and §4. Furthermore, it will lay the groundwork for the results that follow from discussing the relational properties of the models in Chapter §5.

§1.1 Three Historical Models of Reduction:

In what follows I will be canvassing three accounts of reduction. Each is historically significant, and more importantly, each is notably different in approach. First we will look at the often-cited derivational/deductive model by Nagel in §1.1.1. Next in §1.1.2, we will turn to the “disjoint explanation” model given by Kemeny and Oppenheim. §1.1.3 will feature a “model-isomorphic” version inspired by Suppes and developed further by Kenneth Schaffner. I have chosen these models not only because of their prominence in the philosophical literature, but because I think that they represent three distinct conceptual approaches to reduction. Many other models have been proposed, but oftentimes they are similar enough to be grouped with one of the three general archetypes I have chosen.

For each model I will include a discussion about possible candidates for the relata of the reduction relation. Are these only theories? Or are they extended to included theory parts, laws, equation tokens, or models? These questions will become relevant to our discussion of the differing types of reduction that may have been intended by the author, most notable between inter-level reductions and intra-level reductions¹. As future conversations about reduction will illustrate (§3.1, §4.1), there is good reason to think that the domain of the reduction will impose restrictions on the conceptual models that it can be said to exemplify.

Before we examine any of the particulars of the three models, one might wonder how our analysis should proceed. Indeed each author is characterizing what they think reducibility amounts to, yet it is worthwhile to have a discussion about how we might even begin to judge the adequacy of any of the accounts. When deciding what “reduction” amounts to, grounding such a project might initially seem difficult. One way of proceeding might be to examine how scientists employ the word, and seek to give a philosophical treatment of this scientific idiom. A worry here might be that the scientific usages are inconsistent, multifarious, or philosophically lacking². Oftentimes scientists say that there is a “reduction” when they are merely reducing an equation by performing routine mathematical operations after making some assumptions. For instance, I might ask a beginning mathematics student to: “on the Cartesian plane, show that when the

¹ For definitions of many of the important terms in this document, I would refer the reader to the Glossary at §G.

² A similar worry is that scientists, in their daily activity, instead do not use the term “reduction” at all. However, I think that there are a few examples which show that how the term is at least on scientists’ radar, e.g. (Weinberg 1994, chap. 3).

major and minor axes are equal, the equation of an ellipse centered at the origin *reduces* to that of a circle centered at the origin”. This would involve the student looking up the equation for the ellipse in a textbook, imposing conditions on the portions that represent the size of the axes, and then using algebraic rules to transform this into the textbook’s equation for a circle³. Here this does not involve any scientific theories, so it cannot be a reduction. Even if this were situated relative to some orbital problem within a scientific context, it is doubtful it would represent a reduction. Sometimes philosophically-minded goals are not scientific priorities. As a result we might instead forge a conception of reduction that is attentive to philosophical interests. This runs into the problem of seeming *ad hoc*, and is increasingly concerning when a given reductive model is accompanied by cherry-picked scientific examples. Thus, at first glance, when trying to champion a model of reduction, it seems the project is troubled from the onset. In many cases, it would thus appear difficult to provide any positive assessment of a view, that is, to say whether it is *adequate* as a term employed in analysis.

In what follows, I will provide a discussion of each model and its consequences/scope. In reference to the above problem, however, I find the methodology given by Karl-Georg Niebergall to be a helpful guide. He claims that “any attempt to explicate ‘ β is reducible to α ’⁴ can only be acceptable if (i) it is faithful to examples, and (ii) it is a sound representation of general ‘properties’ of reducibility” (2002, 148). I will be providing examples of (i) as each reductive account is presented. The basic foil will be to present a hypothetical/actual case that intuitively seems to represent a reduction, but is one which a given model cannot support. Much of this will occur later as we detail the case studies in §2-§4, and also in §1.2.2. For (ii), our discussion will continue on to §5, where I will provide an extended discussion of the properties that might be traditionally attributed to a reduction.

§1.1.1 Nagel’s Deductive Model:

Nagel was one of the first philosophers to offer an account of reduction; his original article on the subject appeared in 1935. Although he modified portions of his views throughout

³ Symbolically this might involve starting with: $(x/a)^2 + (y/b)^2 = c$, where a, b, c are constants, x, y are variables, assuming that $a = b$, and then attempting to algebraically derive $x^2 + y^2 = r^2$, where r is a constant.

⁴ The symbols employed by this author have been changed. Refer to §G.

the time he spent thinking about the issues, his work generally has a coherent common character. Nagel focused on the syntactic relations that existed between the content (often *qua* symbols) of the reduced theory and the reducing theory. For Nagel, reduction is a derivational matter: if we are able to derive one theory from the other, then we may say that the derived theory is *reduced* by the deriving theory⁵. His account is thus highly logical in character.

Each scientific theory contains terminology that it employs to describe the world – it speaks of atoms, fitness, bonding, or event-horizons. It may be that the reduced theory’s vocabulary is entirely contained by the reducing theories, indicating that there is no new vernacular added by the reduced theory. In this instance we will call the desired reduction *homogeneous*. Any inter-level reduction will seemingly, by definition, contain disparate jargon. Chemistry speaks of different entities than biology, and it measures different quantities with regards to them; this is because, were they to speak of the same entities in the same way, there would be no reason to distinguish the two sciences as different disciplines. Nagel, quite correctly, makes note of this fact, and notices that homogeneous reductions are instead almost always successional affairs that emerge as a theory develops from its predecessors (1935, 339). Relying on a later document (1951), I extract Nagel’s formal account of a homogeneous reduction to be:

β is *homogeneously reduced* to α iff: a derivation of β is possible from α .

Even in naïve intra-level cases, we must be cautious in assessing which vocabulary is shared and which is not. To take an example mentioned by Paul Feyerabend (1962, 80) and Thomas Kuhn (1962, 102), both Newtonian mechanics and relativistic mechanics employ “mass” as a descriptive measure that functions within each theory in a rather similar manner (in that each is an essential factor in describing how an object will move in relation to other objects). However, the objection claims that despite the term “mass” occurring in each theory, it is not in any significant sense the same type of “mass”. Each author claims that Newtonian mass is invariant with respect to an object’s motion, whereas in relativity theory, an object’s mass will increase as the speed of the object increases. So regardless of one theory being historically generated from

⁵ This received historical interpretation of Nagel’s view has recently been challenged (Riel 2011) (Dizadji-Bahmani, Frigg, and Hartmann 2010); we will turn to these challenges later in the section.

another, and in addition to the four-letter word “mass” appearing in each theory’s textbooks, the objection contests that each theory’s use of the locution “mass” is not the same. Inasmuch as we are to claim that we are equivocating two different theoretical constituents, we are to conclude that such a successional reduction could not be deemed homogeneous.

Although the generalized objection seems like a good one, the specific example provided is not. Andrés Rivadulla very convincingly shows that mass *qua* relativistic entity is invariant with respect to measurements at differing inertial reference frames(2004). This gives us compelling reasons to think that the relativistic usage of “mass” is in fact historically contiguous with the “mass” employed by Newton; mass is a “characteristic feature” of objects in the world (Rivadulla 2004, 421). Physics misinterpretations aside, there may perhaps be other legitimate examples which could instead serve the position intended by Feyerabend and Kuhn. For example, usage of the term “gene” by Gregor Mendel and then by modern biochemistry might provide a better case for this difference in meaning⁶.

Another problem recognized by Feyerabend is that the two theories one seeks to compare could be incommensurable (1962, 74–84). For example, one theory might posit that there are two sexes. Another might posit that “sexuality” as a concept is not binary, and that instead there is a spectrum of possible sex identifiers. To attempt a reduction to one theory from another is here hopeless, as no derivation from one to the other could happen, assuming each theory is internally consistent. In some phrases this difference in the locution “sex” does not seem to be a worry, such as with “the sex of a person has an influence on personality”. However this problem is not merely that there is some shared usage of “sex” in each existing vernacular. Instead, the problem would be that there was no basis of comparison or understanding of the new spectrum of sexuality for psychologists who had a strictly binary conception of “sex”. There would be no way for the “male-or-female” camp to ever be able to adequately understand how to place their existing knowledge in terms of this new conception of sexuality. This worry might not be universally damning, as in similar cases inconsistencies arise, yet in these situations we would like to say that a reduction is possible. To take another example, it is well-known that Newtonian dynamics posits space as being Euclidean, whereas relativity theory allows for space to be curved. As the nature of space must be one way or another, a derivation of the former from the

⁶ Schaffner alludes to this as a possible example of the meaning change that Kuhn and Feyerabend had intended, however he notes that this hinges on how we are to employ “meaning” (Schaffner 1967, 138).

latter seemingly cannot exist⁷. However, as Feyerabend sees inconsistencies and instances of incommensurability as being quite common eventualities of theoretical progress, he claims that such a worry will similarly occur with a high frequency for those seeking to apply Nagel's criterion. For Feyerabend, each discussed problem is less the exception and more the rule of theory progression.

When the vocabularies are not entirely shared, a reduction is *heterogeneous*. But if we are to derive one theory from another, we must have a way of making a connection between the exclusive components. Otherwise, short of a contradiction in the reducing theory or a tautology in the reduced theory (a thankful rarity in working science), we will be unable to succeed in completing such a derivation. Thus Nagel employs *bridge principles*, or to use his own idiom, “conditions of connectability” (1979, 354). These are stipulations – often straightforward definitions – that define each of the reducing theory's terms by means of the reduced theory's lexicon. Here the concern is that such bridge principles will be *ad hoc*, constructed without any reasonable scientific foundation. Worse still, *ad hoc* bridge principles might be arranged with the specific goal of assuring that the desired derivation will obtain. For example, a claim germane to particle physics is that “a *kaon* is made from an *up quark* and a *strange antiquark*”. Likewise a claim from chemistry is that “*hydrochloric acid* is made from a *hydrogen atom* and a *chlorine atom*”. Now I may create several bridge principles: “*kaon* ↔ *hydrochloric acid*”, “*up quark* ↔ *hydrogen atom*” and “*strange antiquark* ↔ *chlorine atom*”. Here I have allowed for a Nagelian reduction of a claim in chemistry by a claim in particle physics, although clearly it is not one that anyone would endorse.

To avoid such results, Nagel suggests that a material justification is often required for the supposition of a bridge principle:

For example, one theoretical notion can be made to correspond to the experimental idea of viscosity, and another can be associated with the experimental concept of heat flow. In consequence, since the mean kinetic energy of gas molecules is related, by virtue of the assumptions of the kinetic theory, to these other theoretical notions, a connection may thus be indirectly established between temperature and kinetic energy. Accordingly, in such a context of exposition, it would make good sense to ask whether the temperature of a gas is proportional to the value of the mean kinetic energy of the gas molecules, where this value is calculated in some indirect fashion from experimental data other than that

⁷ We will return to such an example in §2.2, and find that such a worry will not provide any difficulty for such a reduction to occur.

obtained by measuring the temperature of the gas. In this case the postulate would have the status of a physical hypothesis. (Nagel 1979, 356–7)

That some bridge principles may be established scientifically seems readily evident. When we are to decide if there in fact are empirical resources available to establish a bridge principle seems the more difficult question. Likely, it is within the purview of scientists working within each field to make such judgments. For our purposes, the important point to note is that bridge principles are logically necessary to allow for the possibility of a heterogeneous reduction.

As some bridge laws may be empirically established, what sort of claim does this make them? Kenneth Schaffner classifies them as “synthetic sentences” (1967, 138). Recent work has attempted to claim that these may posit ontological associations between entities or properties of either domain (Dizadji-Bahmani, Frigg, and Hartmann 2010, 403–404) (Riel 2011, 364–367). A further worry here is that these ontological claims may in some way be justified by deep-seeded mereological presumptions rather than well-founded *a posteriori* science. The ultimate status of the justificatory mechanisms operating in bridge principles seems difficult to judge except on a case-to-case basis.

Below I will try to provide my own succinct formalization of the Nagel model. In doing so, I rely on later attempts by other authors (Kemeny and Oppenheim 1956, 9–10) (Schaffner 1967, 138) (Riel 2011, 364–367) (Dizadji-Bahmani, Frigg, and Hartmann 2010, 403–404), as well as those of Nagel himself (1979, 356–7). It is important to notice that the Nagel reduction corpus spans over 30 years. Over this time portions of the model have been modified and amended, in reaction to critics and part of a natural progress of Nagel’s ideas about reduction. Indeed when making his own comments leading up to his summary of what the overall coherent “picture” of the Nagel model of reduction would look like, van Riel reminds us that “there is no such picture” (2011, 364–367). Here is my construction of the Nagel model:

β is *heterogeneously reduced* to α iff:

- (I) The theoretical vocabulary of β contains terms not in the theoretical vocabulary of α .
- (II) Each term of β is linked to a term or compositions of terms in α by means of well-established biconditionals.
- (III) The derivation of β by means of these biconditionals follows from α .

One notable move in the above is how I have chosen to represent the bridge laws that form **(II)**. Some of what is said in Nagel makes it appear as though the biconditionals are to link singular terms of β to *singular* terms α (for example [Nagel 1979, 356–7]). However, upon reflection one can imagine many cases in which we would want multiple components of α to combine to make some component of β . Indeed, anytime there is a mereological relationship in which several different parts of α composed β , such an allowance would be essential. Allowing a link to “terms” of α , as opposed to “a term”, is also consistent with other exegeses (Schaffner 1967, 138). Finally, when considering Nagel’s own example of linking “temperature” with “mean kinetic energy of molecules” (1979, 356–7), I see the latter locution to be difficult to formalize into first-order logic without the use of multiple predicates.

In one tantalizing footnote (Nagel 1979, 356–7), Nagel briefly considers that we might have the conditions of connectability that are described by **(II)** be conditionals instead of biconditionals⁸. This left some authors to hypothesize about what hangs on such a change (Bickle 1998, 120) (Richardson 2008, 403). However, as Ronald Endicott notices, directly after the suggestion in the same footnote Nagel rescinds the suggestion, on the grounds that it would not allow for α to replace β (R. P. Endicott 2007, 3–4). As such, I feel that the representation best captures Nagel by keeping the biconditionals.

Worries about the existence of inconsistencies prompt a few comments. First, we could show that an inconsistent theory reduces *any* theory, as a consequence of its ability to derive anything. To assuage this problem, we could simply require that candidates for the α relatum be “legitimate” or “seriously considerable” scientific bodies⁹. Similarly, if we were to employ two separate theory parts, α_1 and α_2 , to reduce a β , it may well be that $(\alpha_1 \ \& \ \alpha_2)$ is inconsistent despite α_1 , α_2 both being consistent on their own. The fix here is to simply stipulate that the constituents of the reducing theory set, when taken as a whole, must meet the “legitimacy” requirement. Finally there is a concern about having a perfectly consistent α that, when combined with the biconditionals from the connectability requirements, entailed a contradiction.

⁸ This would thus amend **(II)** to read as follows:

(II)* Each term of β is the consequent of a well-established conditional whose antecedent is comprised of a term or compositions of terms in α .

⁹ Similar maneuvers will be made by later authors to eliminate the worry of fanciful examples: Nickles does so in §3.3 by his “establishment” provision, while Rohrlich and Hardin do so in §3.4 via their “maturity” clause.

However I think we would avoid difficult cases such as these, once we have imposed the joint requirements of “legitimacy” for α , and “well-establishment” from (II).

We must be cautious when comparing Nagel to other thinkers. Nagel was one of the first to write on reduction in the sciences, and since that time much has changed in the philosophy of science. To provide an example, Schaffner claims that a reduction is nothing more than one theory explaining another (2006, 380). Nagel himself seems to corroborate such a point as he claims that “reduction, in the sense in which the word is here employed, is the explanation of a theory or a set of experimental laws established in one area of inquiry, by a theory usually though not invariably formulated for some other domain” (1979, 338). We must tread very carefully with quotes like these. Without keeping in mind the context of argument, the two authors might be said to agree. Instead I think it is correct to claim that Nagel’s criterion is an “explanation” only in the sense that the word was understood as being an application of the deductive-nomological (D-N) conception of explanation, as developed by C. G. Hempel (1970). We can imagine Nagel disagreeing with Schaffner if he knew what broad usage of “explanation” the quote allowed. Furthermore, we could then rely on the rather significant literature criticizing the D-N model of explanation¹⁰ to show how, in the respect that Nagel intended, there are reductions that do not explain *qua* D-N explanation.

The received view of Nagel’s position has traditionally attributed it to be a relation between two theories (Hoyningen-Huene and Wuketits 1989, 30). For example, my above extraction speaks only of α and β as singular theories. Also, most of Nagel’s language when speaking on reduction concerns “theories”, as he more generally wanted his discussion of reduction to be applicable to reductions that range over several theories. Nagel’s diction is especially troubling when we consider that he has a rather austere notion of what a “theory” consists in (1979, chap. 11). Under Nagel’s definition, it is difficult to see how psychology or even areas of biology could be considered theories. Hence reduction in these areas, for many the frontier cases for active inter-level science, simply would not work.

Through a very broad textual analysis, van Riel argues to the contrary. He claims that Nagel conceived his model in relation only to ideal cases, where theories are explicitly provided in terms of axioms (Riel 2011, 360–362). Furthermore, when looking at later cases where Nagel discusses the possibilities of reduction for cases of biology, and also for property reductions or

¹⁰ For an excellent historical account of this criticism I would refer the reader to (Salmon 2006, 46–50).

reductions that obtain between parts of theories, van Riel notices that Nagel is in each case optimistic. This leads van Riel to reason that, rather than merely being a relation between theories, we should instead read Nagel as believing that “reduction is a relation holding among a great variety of scientific representational devices, among which theories play an important epistemological role” (2011, 362).

To conclude this discussion, I think that Nagel rightly recognized the possibility of reduction between sorts of entities that might not be considered “theories” under most definitions. He admits that he is exclusively working with idealized theories, and that these may not track actual scientific practice at a given time (Nagel 1979, chap. 11). So van Riel is correct in his Nagel exegesis. However the focus of many who wrote on Nagel, this section included, was on his model of reduction, and this was indeed intended as something that allowed only theories. Thus I think that many authors whom van Riel cites as maintaining that the Nagel model’s relata may only be theories (for example [Sarkar 1992]) are also correct. There is an important difference between what we attribute to *Nagel on reduction*, and what we attribute to *Nagel’s model of reduction*¹¹. The latter was the focus of Kemeny and Oppenheim, of Sahotra Sarkar, and also of nearly every other author that has discussed Nagel. Thus I will conclude that Nagel’s focus, and certainly what he allowed with his reduction model, concerned “theories”.

Traditionally, Nagel’s model has been classified as *direct*, in that it deals merely with the two theories and does not make any essential reference to information beyond what is stated by those theories (Schaffner 1967, 138). So long as we restrict ourselves to Nagel’s *model*, I think that such a classification is accurate: the model concerns the axiomatic representation of each theory and a collection of biconditionals whose constituents are terms from each theory. As every part of the process – the derivation and the parts employed with the derivation – is just a part of either theory, we can conclude Nagel’s model is a paradigmatic direct reduction. There have been worries about the essentiality of characterizing theories in first-order logic (Dizadji-Bahmani, Frigg, and Hartmann 2010, 400), as well as general problems concerning the possibilities of an indirect/direct dichotomy for theory explanation (Riel 2011, 367–370). However as our focus will be on the model, I feel these issues may be ignored.

¹¹ This is certainly an issue that van Riel is aware of, for his paper was titled *Nagelian Reduction Beyond the Nagel Model* (2011). My point in this paragraph is to illustrate the difference between the two possibilities, and also let the reader know where our focus will generally lie.

To summarize the important features of the Nagel model, it provides logical means for there to be a direct derivation of β by α . The non-mutual terminology is connected by biconditional bridge principles that are then included in the derivation. Importantly, since all that is added are connecting principles, β is effectively contained within α . All that may be posited from β , its predictions, explanations, etc., must also be able to be posited from α . It requires exact correlation between the two theories, and leaves no room for error in β . This is problematic when we seek to fit it to actual scientific cases, as we will elaborate in §1.2.2 as well as in §2.1. Indeed, we will notice how the cases studies from §2 and §3 both fail to fit the Nagel model due to this stricture.

§1.1.2 Kemeny and Oppenheim's Disjoint-Explanation

Model:

After the Nagel model, Kemeny and Oppenheim proposed a different model of reduction. The Nagel model asserts that a reduction obtained when there is a sufficient connection between the two theories, and that furthermore this connection is to be derivational. Kemeny and Oppenheim instead saw that the hallmark of a successful theory is that it was good at accounting for what happened in the world, that it could give an explanation of empirical data. Were it to be that one theory did a better job at explaining the same data as another theory did, then we should prefer this more successful theory. The same notion can be applied to reduction: one theory is reduced to another when it can explain data at least as well, and also be better organized, presented, and conceptually contained.

Unlike Nagel, Kemeny and Oppenheim succinctly provided the formal conditions that would satisfy their model (1956, 13):

Red(α , β , O) [a reduction of β by α relative to observational data O], if:

- (I) The theoretical vocabulary of β contains terms not in the theoretical vocabulary of α .
- (II) Any part of O explainable by means of β is explainable by α .
- (III) α is at least as well-systematized as β .

Although a feature that has warranted little discussion, **(I)** was perhaps put in place as homage to Nagel's "heterogeneous reduction". It also has some consequences for the entities that will qualify as being reduced under the model, as we will later elaborate.

The most significant novel feature of the model is its explicit reference to the empirical, by its consideration of the data set **O**. As perhaps expected, a significant question arises about what a theory "explaining the observational data" amounts to. There has been considerable discussion of scientific explanation since the 1955 paper. For example, Bas van Fraassen contests that an explanation is a satisfactory answer to a "why" question, one that is unavoidably context dependent (Fraassen 1980, 156). Under this account of explanation, so long as the associated "why" questions inspired by **O** that are adequately answered (relative to the given context) by β are likewise answered by α , then this is sufficient to satisfy **(II)**. Were we to presume a different model of scientific explanation, the status of the reduction has the potential to change. Thus not only is the satisfaction of **(II)** dependent upon **O**, it also hinges on what the correct account of explanation is thought to be. The matter is further complicated when we consider that this later observation presumes a monism for explanation in the sciences. Considering the good number of philosophers who are content to admit a plurality of models of explanation (Kellert, Longino, and Waters 2006), we may find that a given β and **O** might reduce to many disparate α 's. This has the potential to disrupt several entrenched notions for the unity of science, such as the uniqueness of successors in successional reductions, or the pyramidal structure of inter-level reductions¹².

Current interpretations aside, at the time of their reduction-model's conception Kemeny and Oppenheim believed that their model was consistent with "any explicatum of 'explain'" (1956, footnote 18). Thus employing van Fraassen's account with **(II)** would be warranted in some cases, despite the anachronism. However, work on scientific explanation had not developed as much in 1956 as it has today, and there were few other contenders to the D-N account, considered at the time to be the received view. In their paper Kemeny and Oppenheim employ the D-N account to prove a theorem (1956, footnote 18), and much of their discussion reads as if the authors have the D-N account operating behind the scenes all along. This should be unsurprising: one of the earlier papers written by Hempel in 1948 was co-authored by

¹² These two "sacred cows" of classical reductions are discussed, along with other possibilities, throughout §5.2.

Oppenheim. Indeed, even today, many prefer to refer to the D-N model as the “Hempel-Oppenheim model” (Kitcher and Salmon 2011, 38).

Under the D-N theory, general “covering laws” are combined with particular statements endemic to the phenomenon in question to deduce the *explanandum*. Oftentimes this implies that “certain quantitative features of a phenomenon are explained by mathematical derivation from covering general laws” (Hempel 1966, 52). Say that a covering law, together with circumstantial data, predicts that an eclipse will occur at 21:00. If we find that the eclipse in fact occurred at 21:01, did the theory account for (viz. explain) the eclipse adequately? What if the eclipse had occurred at 21:13? At this time it seems that the correct answer is that such worries are not a problem for the reductive criterion suggested by Kemeny and Oppenheim, but instead an interesting conundrum for the D-N model of explanation. I mention them here because such issues concerning the accuracy of correlation, although not important now, will become problematic as we delve further into Kemeny and Oppenheim’s account.

Another interesting issue that arises under this reductive model concerns the status of false theories in reductions. If we claim that Ptolemaic astronomy is reducible to Newtonian mechanics, or that Newtonian mechanics is reducible to relativistic theory, then we must at some point consider whether we should to allow false theories to explain observations. For many, Ptolemaic astronomy can predict certain aspects of the observational data quite readily. However it is still a false theory, in that the account of the motion of astronomical bodies it provides is incorrect, and not by just a little. Thus many see it as being impotent at the explanatory level¹³. If we are not going to allow Ptolemaic astronomy to explain any of the data, then satisfying (II) either becomes impossible or automatic, depending on how one interprets the statement. Given the requirements of (II), we may find ourselves in a very curious position when attempting to speak of reductions that involve admittedly false theories. This worry is especially an issue when dealing with successional reductions, as in almost all cases the falsity/ineffectiveness of the succeeded theory is granted¹⁴.

The authors also spent little time explicating how “systematization” was to be read in (III). It was understood to be a “measure that combines strength and simplicity, in which

¹³ Not all are willing to admit that this is in general a problem for so-called “fictitious” theories. Recently Alisa Bokulich has tried to argue that in certain cases false theories may indeed be found to adequately explain observations (Suarez 2008, chap. 6).

¹⁴ Later we will revisit the impact that false theories have on successional reductions in §2.1 and §3.1.

additional complexity is balanced by additional strength” (Kemeny and Oppenheim 1956, 11). Thus when choosing between a first-order logic with a left-identity operator and a right-identity operator, and a first-order logic with a single, “both ways” identity operator, we would choose the latter on grounds of simplicity, despite the former being just as adequate in proving theorems. Similarly we would prefer a system with identity to one without, on the grounds that the simplicity (as measured by number of symbols and rules) of the system without identity comes at a sacrifice of power, as exhibited by the semantic fecundity of the identity symbol for translations and derivations. Kemeny and Oppenheim see “systematization” as intuitive, while admitting that “the concept is in need of precise definition” (Kemeny and Oppenheim 1956, 11).

Kemeny and Oppenheim rightly recognize that the body of accepted observational data quickly changes. They thus time-stamp the observational data as ${}^t\mathbf{O}$. However Kemeny and Oppenheim attempt to be rid of the “undesired parameter” that is $\mathbf{O}$. They believe that by considering every possible data set that would be amenable to a β -explanation, if such sets could also be explained by α , it would seem that the particulars of the current body of observational knowledge could be eliminated. The authors therefore suggest the following definition:

Definition 6: **Red(α , β)** if for every $\mathbf{O}$ consistent with β , **Red(α , β , $\mathbf{O}$)**

Again, it seems noteworthy that such a definition is contingent upon what “consistency” amounts to in this case. In what follows, I will use the D-N account when doing any explaining, as this was the preferred model of Kemeny and Oppenheim. First let our sample β be a very simple theory whose observational predictions are dictated by a single function that relates two variables, $\mathbf{r}$ and $\mathbf{s}$:

$$\mathbf{s} = 5\mathbf{r} \quad [1.1]$$

Now look at the following observation sets, representative of observational data, organized as ordered pairs $(\mathbf{s}, \mathbf{r})$:

$$\mathbf{O}_1 = \{(2, 10)\}$$

$$\mathbf{O}_2 = \{(2, 10), (3, 15), (10, 50)\}$$

$$\mathbf{O}_3 = \{(5, 25)\}$$

Here it seems safe to say that $\mathbf{O}_1$ is “consistent with β ”, as all of the values directly correspond to those that β would have predicted. Now if we are to have $\text{Red}(\alpha, \beta, \mathbf{O}_1)$, recall that every part of $\mathbf{O}_1$ that is explainable by β must also be explainable by α . Let us presume that all of the required factors beyond the quantitative correlation obtain in the right sorts of ways between our β and $\mathbf{O}_1$, so that we can confidently claim that β explains $\mathbf{O}_1$. These factors may be quite complicated or rather minimal, depending on which account of scientific explanation is employed. Relying on the D-N model, it is enough to notice the correlation between $\mathbf{O}_1$ and the values arrived at by entering differing inputs to the “law” that is [1.1]. $\mathbf{O}_2$ is also consistent with β , and additionally $\mathbf{O}_2$ contains more data than $\mathbf{O}_1$, as $\mathbf{O}_1 \subseteq \mathbf{O}_2$. $\mathbf{O}_3$ just represents another possible data set that is also unquestionably consistent with β .

Kemeny and Oppenheim are quite explicit about why they felt the need to invoke the “consistency” clause that is *Definition 6*. The worry is that there may well be situations where β and $\mathbf{O}$ are inconsistent, thereby entailing any proposition. One related problem is that β itself is inconsistent, before any introduction of $\mathbf{O}$ ’s. Rather than specifying to eliminate this worry through definitions, Kemeny and Oppenheim were willing to admit that a reduction can occur in this case, albeit ones that are rather “uninteresting” (1956, footnote 11).

Keeping in mind the requirement for consistency to β , *Definition 6* forces one to first consider the range of observations sets that can be explained by β . This in our example would become:

$$\text{The power set of } \{(x, 5x) \mid x \in \mathbf{R}\} \quad [1.2]$$

Next each of these members must be checked to show that α could also explain them. I would like to note that even with our simple example such an endeavor is potentially arduous. There are an uncountably infinite number of members of [1.2], and some of these members are themselves uncountable infinite sets that do not have any well-described progression. Notice additionally that [1.2] is an undecidable set, and that a good portion of its members are themselves undecidable sets. For instance take a look at this set:

$$\{(x, 5x) \mid x = k10^{-n}, \text{ where } k \text{ is the } n\text{th digit of } \pi\} \quad [1.3]$$

Here $[1.3] \in [1.2]$, $[1.2]$ can be succinctly described, and furthermore $[1.3]$ is even decidable. However such a set has the potential to be quite cumbersome to work with, especially when the laws involved in α and β might require some mathematical cleverness to divulge all potential data points.

Other questions arise when considering how much leeway we are willing to grant in interpreting the universal claim that is *Definition 6*. If we find that α may account for every possible member of $[1.2]$ except for:

$$O_4 = \{(105, 210), (250, 500)\}$$

we could certainly just hold fast and admit that in this circumstance we had not satisfied *Definition 6*'s requirements, and thereby claim that **Red**(α , β) did not occur. However I hope that for some there is an intuition that, with just one member missing from the uncountably infinite set that is $[1.3]$, to disallow that **Red**(α , β) is a bit draconian.

It may well be that the last few points are not especially relevant scientifically, as we rarely employ uncountably infinite data sets. I will leave that issue aside to turn to another more significant worry. Consider:

$$O_5 = \{(1.1, 5.07), (1.98, 10.011), (2.87, 14.99)\}$$

The above is not, strictly speaking, the output of $[1.1]$. Yet still we must wonder: does this mean that this set is compatible with β or not? Even if we are to grant that it is, it is another matter to address whether such a set is indeed explainable by β . Earlier in the paper Kemeny and Oppenheim state that they “intentionally overlook the fact that most observations involve a margin of error” (1956, 8). I read this statement as an attempt to curtail worries about “how close” a theory’s predictions need to be to the observed data. This is similar to the situation discussed above where the D-N model of explanation was forced to consider observational results that deviated from the theoretical results. Recall that before I reasoned that this was more a problem for the D-N theory of explanation to resolve, as well as possibly being a question that

concerns the process of theory confirmation. However I do not think that Kemeny and Oppenheim in the current circumstances are permitted to presume that such an issue can be likewise bracketed. Instead I will argue that, as a consequence of *Definition 6*, we must address this issue of theory/data deviance and all its nuances.

These details may seem overly picky, but I am interested in how one would go about checking data sets in the specific. I think that when we look at specific cases, problems come up that weren't otherwise apparent. For example, under any modest philosophy of theory confirmation and explanation, the story of whether the [1.1] of β fits and ultimately explains the data set O_5 will involve a contextual analysis of the circumstances of the experiments that yielded that data. When were they taken? How many samples? What was the accuracy of the equipment? Has such an experiment been conducted before with similar results? Are there any experimental reasons that would cause results to be notably skewed/varied? These questions and others are almost certainly likely to be named as contributing factors in the decision as to whether or not the theory can be said to explain the observational data – this is how scientists determine what constitutes “good data” and what does not. As such, I can vary these contextual features so that β may be found to very adequately explain O_5 . However by changing these experimental circumstances I can make it so that the same β will be found to not predict/explain/account for O_5 . This may be done all while additionally holding our model of scientific explanation fixed. As a mere grouping of numbers, if an observational set differs even slightly from the set calculated from theory, we are unable to pass judgment as to its adequacy. To say that it either “explains” or “doesn't explain” an imperfect set such as O_5 requires that one look at details about said observational set's generation, among other things; suffice it to say that such an investigation must include factors beyond a mere statement of data points.

Granting this, universalizing about O becomes unbelievably difficult. Looking at the range of data sets that β might be said to explain, we are left with [1.2], *in addition to* the more complicated consideration of the deviant data sets such as O_5 . Another example of a deviant data set that would be up for consideration could be a modest finite subset of [1.3], but with each of the elements modified by $\pm .01$ ¹⁵. So *each* of the deviant sets must be additionally accompanied by an account of the experimental circumstances under which they were produced if we are to

¹⁵ I mention this additional example to show that, in permuting it and making it finite, this set is now less fanciful. Instead, it is potentially patterned closely to what we might find in an actual real-world observation set.

determine whether or not they are explained by β . Thus, *Definition 6* would require that we produce:

$$[1.2] \cup \{(S, C_s) \mid S \in \text{the power set of } \{(x,y) \mid x,y \in R\}\},$$

Where C_s is any circumstance where β would explain S

Every element of this monstrous set would then have to be checked to see if α would be found to likewise explain it! The number of cases we could construct that would allow just O_5 to be explained by β is large, and varies wildly. This is akin to the task of considering all of the circumstances that I would believe my brother when he told me that he was in a good mood – to attempt to classify and enumerate the various situations is a herculean task. Each of these cases would have to be formulated within a range of specifics and qualifications. For instance: “my brother is in a good mood” if he has a surprise party thrown for him and his friends are all there: but not the ones he dislikes; no one in the family has recently died (that he liked); the party goes well (there are no fights, tragedies, etc.); he remains in good health; etc. Perhaps many of the provisions may fall under the purview of a broader *ceteris paribus* assumption, but in some cases I think that this dismissive hand-waving occurs too quickly. For actual scientific examples, we can imagine all of the subtleties that are involved when expert practitioners participate in taking data. If we allow that the bacteria of a sample must be observed under a “microscope that is clean”, this condition alone is full of minutiae that will speak of those cleaning it, how it was cleaned, which state of repair the microscope may be allowed to be in, which types of microscopes are permissible, etc. Concepts like these necessarily cling to each datum that became involved in the verification of a theory. However we often rely instead on the testimony of the practitioners who have evaluated the relevant data, thereby qualifying it to be in the range of often-unspoken (and frequently contentious) conditions. But were we to consider all possible observational instances, we presumably cannot circumvent a consideration of even mild border cases.

All of these concerns are still background to the final task of deciding if, in each case that we admit that a given O was explained by β , this O would also in this case be explainable by α . For were α and β to differ in the slightest in their theoretical posits, equations, explanatory demands, etc., then we are likely to find this final task of determining that α does in fact explain

an **O** to be very different than the decision making process employed when such an **O** was first found to be explainable by **β**. In these cases especially we would think that a few well-placed *ceteris paribus* qualifications wouldn't help. For the cases where a competent biologist were to decide that a microscope was clean might not be able to compare in any helpful way to the manner that a competent chemist would require for a beaker to be uncontaminated.

To conclude, I think that Kemeny and Oppenheim are likely justified in ignoring theory/observation interface issues when stipulating their reductive criterion for **Red(α, β, O)**. This is because in this context the problem of whether a theory explains an observation that deviates slightly may really be bracketed as an issue for the philosophy of theory confirmation/acceptance, as well as of course depending on one's preferred model of scientific explanation. However by attempting to eliminate **O** by use of their *Definition 6*, Kemeny and Oppenheim ignore the crucial and complicated role that the experimental circumstances in which a given **O** is generated play when deciding if a theory may be said to explain such an **O**; again this is by and large independent of how one might seek to characterize such a process of explanation. Thus I conclude that reference to observations cannot be done away with as easily as Kemeny and Oppenheim have presumed: *Definition 6* is inadequate, because any attempt to demark all possible sets of observations **O** that **β** may be said to explain will become hopelessly bogged down by the magnitude of such a task before a consideration of **α** may even begin to be entertained.

Kemeny and Oppenheim see reduction as a three-place relation involving the two theories and also the available observational data: **Red(α, β, O)**. The above reductive definition given by Kemeny and Oppenheim makes explicit reference to observational data in their schema. Recall how Nagel's account features a translational component. Each piece of the disjoint vocabulary of each theory is then interfaced by means of bridge laws. The rationale for such a requirement is that by permitting the two theories to "talk to one another", one might then be able to mathematically/logically derive **β** from **α**. Kemeny and Oppenheim resist such maneuvers in the above definition. Instead all that is required is that **α** be able to explain the observations that **β** did. It is thus possible that there need not be *any* significant similarities in the explanatory components of each theory; it could very well be that each theory posits entirely different entities, theoretical assumptions, or metaphysical worldviews – they could be inconsistent, conflicting, or incommensurable in the strongest Kuhnian/Feyerabendian sense. Furthermore we

might have little-to-no idea how to translate constituents of one theory into constituents of another. A reduction can occur so long as each theory can reasonably explain the observational data: to account in some way for what the result was, and to answer a few other questions about how the situation surrounding the data resulted in the observation, why the result was not otherwise, etc. The details of this process will depend upon one's preferred model of explanation. Thus the typical classification is to label Nagel's model a *direct*, and Kemeny and Oppenheim's model *indirect* (Schaffner 1967).

Kemeny and Oppenheim see one of the novel differences from their model to Nagel's is the involvement of the empirical: "the essence of reduction cannot be understood by comparing only the two theories; we must bring in observations" (Kemeny and Oppenheim 1956, 13). The mechanism of reduction employed by Kemeny and Oppenheim is certainly different, in that it only relies on links between the theories and observations. But recall that a heterogeneous Nagelian reduction required bridge laws, the viability of which depended upon observational evidence. So, inasmuch as we understand a reduction to be providing an explanation of one theory by another, we cannot claim this is done solely by the resources provided by α alone. Instead both α and the bridge laws are components of the *explananda*. This component was not acknowledged by Kemeny and Oppenheim, but it is essentially employed by each of the two models. Thus it is wrong to claim that Nagel's model does not involve observations: the model merely neglects to make explicit reference to them.

Kemeny and Oppenheim choose to have theories expressed in a "formalized language". This is less a comment about how science actually behaves than an aid for their logical treatment of the issue. However it is clear that they envision theories to consist of laws, equations, and generalizations, and are interested ultimately in reducing "branches" of science to other branches. There is little in their discussion to suggest that α and β are anything more than syntactic groups of axiomatic expressions in the traditional sense.

The authors make no distinction between inter-level reductions and intra-level reductions¹⁶, but refer to a special case of reduction where the vocabulary of α is contained in β , calling this an "internal reduction". Understanding $\text{Voc}(\delta)$ to mean *the vocabulary of theory δ* , they state:

¹⁶ Later Oppenheim and Hillary Putnam define "micro-reductions", used to primarily demark inter-level reductions (Oppenheim and Putnam 1958, 7). I have chosen to leave this distinction unexamined, as I feel that further compositional requirements add little to the discussion I have made in this section.

*Definition 5: **Intred**(α , β) [an internal reduction of β by α], if and only if:*

Red(α , β) and **Voc**(α) is a proper subset of **Voc**(β)

(Kemeny and Oppenheim 1956, 14)

This does not easily map onto our concepts of inter- or intra-level reduction. First we can envision an intra-level reduction that isn't an internal reduction, such as the case of ray optics being succeeded by wave optics: there are terms in the wave optical lexicon that are not employed in ray optics, such as "interference". Likewise we can conceive of an intra-level reduction that is an internal reduction by using the example of Newtonian mechanics and special relativity. Each uses the same terminology, "mass", "velocity", "speed of light", etc., yet as James Weatherall, following many others, has noticed, there is a difference between how each theory considers "mass". Newtonian physics has an "inertial mass" and "gravitational mass", whereas relativistic physics merely has one "mass" concept (Weatherall 2011)¹⁷. On the other end, every case of an inter-level reduction I can imagine isn't internal. This is because the "more fundamental" theory seemingly must refer to broader, composite entities that are not part of β 's ontology. So, if there is any work done by the concept of an "internal reduction" in relation to the inter-level/intra-level difference, it is that every internal reduction is intra-level. However not every intra-level reduction is internal.

Whether we are employing the internal reduction schema or the more general reduction model, we are never allowed to have **Voc**(α) = **Voc**(β). Yet it is not much of a stretch to imagine an intra-level reduction where all the terms being employed were identical, yet the way they functioned in the theories/equations was changed. Each theory here would describe the world differently, not with regard to the vocabulary or entities, but in how they were interrelated. The focus isn't terminological but structural. This case could satisfy (II), as the explanations offered by β could simply be less satisfying than α 's, and likewise (III) could happen. However, this would be disallowed as a reduction because of (I). So at the cost of mirroring Nagel's "heterogeneous" criterion, it seems that Kemeny and Oppenheim's model is unable to capture some cases that seem like excellent candidates for a reduction.

¹⁷More on this in §2.2.1.

In summary, the Kemeny and Oppenheim account of reduction is an indirect model that sees the principle function of theories as providing explanations of the empirical. Also, the preferred theories are ones that are more systematized. We may say that at a given time t , relative to a data set O , there are explanations of this data by α and β , or we may eliminate these qualifiers by showing that for every O , it is the case that α provides superior explanations to β . A significant problem occurs when examining what such elimination would entail, for the nuances involved with stipulating what “explaining O by a theory” amounts to for *every* O result in an overly onerous requirement. Finally, the model is indirect in that there never needs to be any interface between the vocabulary, concepts, etc. of α and β .

§1.1.3 Suppes’s Semantic-Isomorphism Model:

In his foundational *Approaches to Reduction* (1967), Schaffner brings attention to an account of reduction that can be formulated from two brief remarks made by Suppes (1957, 10:271) (1967, 59). The account seemingly avoids any direct relations between two theories and instead seeks similarities in their consequences. However, as opposed to considering what the two explain, the idea is instead to focus upon what they semantically *entail*. Assume, as usual, that we are looking to show what a reduction from β to α would amount to. Here the requirement will be that:

For any model of β , there exists a model of α such that the two models are isomorphic. [1.4]

Now the task is to show exactly what constitutes a model of a theory, and then to show how an isomorphism between these two entities will be said to obtain.

It should be noted that, when Suppes made his remarks in 1957 and 1967, a rigorous semantic view of scientific theories was still in development. Indeed Suppes was one of the early pioneers of this movement (1967). Suppes was interested in a semantic modeling of theories once they had been axiomatized (that is to say, rigorously treated in logical or set-theoretic terminology). For first-order logic, a *model* M of a theory T is a non-empty universe U and an interpretation function I that:

1. Maps all of the $\mathbf{n}$ -place predicates $\mathbf{P}^n$ of $\mathbf{T}$ to sets of $\mathbf{n}$ -tuples whose entries are elements of $\mathbf{U}$.
2. Maps all of the $\mathbf{m}$ -place functions $\mathbf{f}^m$ in $\mathbf{T}$ to elements of $\mathbf{U}$.
3. Maps every constant in $\mathbf{T}$ to an element of $\mathbf{U}$.¹⁸

We additionally require that $\mathbf{M}$ will make $\mathbf{T}$ true, that is $\mathbf{M} \models \mathbf{T}$. In relation to his comments on reduction in 1957, Suppes remarked that “a satisfactory general definition of isomorphism for two set-theoretical entities of any kind is difficult if not impossible to formulate” (1957, 10:262). Subsequent writers have disagreed with this claim, and adequate notions of isomorphism have been given from many sources with little controversy. Taking my cue from (Jech 1978, 155), two models $\mathbf{M}_1$ and $\mathbf{M}_2$ are *isomorphic* ($\mathbf{M}_1 \approx \mathbf{M}_2$) when there is a bijective function ϕ that will map $\mathbf{U}_1$ onto $\mathbf{U}_2$ so that the following hold:

- i. For all $\mathbf{n}$ -place predicates $\mathbf{P}^n$, $\mathbf{P}^n_1(\mathbf{x}_1, \dots, \mathbf{x}_n)$ if and only if $\mathbf{P}^n_2(\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_n))$.
- ii. For all $\mathbf{m}$ -place functions $\mathbf{f}^m$, $\mathbf{f}^m_1(\mathbf{x}_1, \dots, \mathbf{x}_m)$ if and only if $\mathbf{f}^m_2(\phi(\mathbf{x}_1), \dots, \phi(\mathbf{x}_m))$.
- iii. For all constants $\mathbf{c}$, $\phi(\mathbf{c}_1) = \mathbf{c}_2$.

It follows from the above that two isomorphic models will in all cases possess coextensive truth evaluations. This is to say that, granting $\mathbf{M}_1 \approx \mathbf{M}_2$, for any sentence θ , $\mathbf{M}_1 \models \theta$ whenever $\mathbf{M}_2 \models \theta$. Now with the above definitions in mind, we can restate [1.4]:

(I) For every $\mathbf{M} \models \beta$, there exists an $\mathbf{M}' \models \alpha$ such that $\mathbf{M} \approx \mathbf{M}'$.

Suppes’s account of reduction is quite interesting, as the relation between the two theories is entirely structural. As nearly every author who has worked with such an account has noted, (I) allows for cases where a reduction occurs between two theories that have seemingly little in common (vocabulary, subject matter, etc.), but happen to possess similar structures (Schaffner 1967, 138) (Bickle 1993, 366) (Sarkar 1992, 169). This is problematic as it allows psychology to reduce to physics, in principle, so long as the theories do not share any collective

¹⁸ Here I rely heavily on the work of Thomas Jech in formulating these definitions (1978, 155).

vocabulary between the two of them. One answer to this trouble might be to note that more is needed to capture the entire relation, as each of the above authors suggested that **(I)** is a necessary condition to reduction, but not sufficient.

The model-isomorphism strategy that Suppes inspired was first expanded by E. Adams (1959), and has been resuscitated in recent years as a “structuralist account of reduction” (Balzer, Moulines, and Sneed 1987)¹⁹. Rather than give the technical details of each model, I will take help from Thomas Mormann. He correctly concludes that despite the various structuralist accounts differing in detail – progressively becoming more resolute – they have a common character: “a structuralist reduction relation between a reducing theory α and a reduced theory β is a structure-preserving map between the structured objects α and β ” (1988, 217). The map in each case is merely some morphism between the objects in question, where the objects in question are highly-mathematized objects that are interrelated in a specific manner. The structuralist elaborations of this version of reduction continue to be developed and criticized, although I will not pursue them much hereafter²⁰.

There is a worry about how this archetype can deal with certain improvements on the predictions of a theory. Take a simple example, where β claims that an object is in motion just in case it has a non-zero net force acting on it. Also assume that β posits the existence of massless objects, and has assumed the motion law does not in any way change for the physics of things without mass. Later, theorists improve upon the theory, by being able to consider the dynamics of massless objects. They have found that the previous theory’s law (stating that a non-zero net force is necessary and sufficient to be put in motion) is a perfectly good law for objects that have mass, but that it does not apply to massless objects. To play by the rules that we have sketched-out above, I will formalize our example in first-order logic. Here is our schema of translation:

F τ : the net force acting on τ is non-zero
G τ : τ is put in motion
H τ : τ is massless

This allows us to formulate β as:

¹⁹ This is of course a very fast summary. I would direct the reader to Niebergall (2002) or Mormann (1988) for a more adequate catalogue of the structuralist literature.

²⁰ I mention most notably the seminal criticisms provided by Niebergall (2002) of each of the major structuralist conceptions of reduction. Indeed, his method of juxtaposing grounded pre-theoretic intuitions against analyzed theoretical results, which I mentioned in §1.1 was an inspiration for a portion of my own methodology in analyzing models of reduction.

$$(x)(Gx \leftrightarrow Fx) \ \& \ (\exists x)Hx$$

Likewise translating we will find that α may be represented as:

$$(x)((Gx \leftrightarrow Fx) \leftrightarrow \sim Hx)$$

Now the following model, call it $\mathbf{M}$, will satisfy β :²¹

$$U=\{1, 2\} \quad F:\{1\} \quad G:\{1\} \quad H:\{1\}$$

I certainly could have made a model that was perhaps closer to real-world cases, with thousands of objects that have F , G and just a few that lack H . But this is not necessary: $\mathbf{M}$ is a possible world and $\mathbf{M} \models \beta$.

Now let us consider all of the possible models that satisfy α . $\mathbf{M}$ will not, but that isn't necessarily a problem. This is because when making an isomorphism, rule **ii** above allows us to choose any function ϕ that will make a one-to-one map from a predicates assignment from β onto the predicates in α . This means that we may choose to pair the F of β , hereafter called F^β , with any of the predicates in α : F^α , G^α , or H^α .²² Thankfully I have chosen $\mathbf{M}$ so that these options will not matter. Each of these choices will generate the classes of models, all of which will be isomorphic to the following model, $\mathbf{M}^*$:

$$U^*=\{1, 2\} \quad F^*:\{1\} \quad G^*:\{1\} \quad H^*:\{1\}$$

And it should be apparent that $\mathbf{M}^*$ cannot satisfy α . This means that there is no model that satisfies α that is isomorphic to $\mathbf{M}$, as every possible model in α will be isomorphic to $\mathbf{M}^*$. So **(I)** will not be satisfied and there cannot be a reduction under the Suppes archetype.

²¹ Proof of this fact, as well as a direct demonstration of the following isomorphisms, are left to the reader.

²² One might worry about an unrestricted cross-mapping of predicates, as it can possibly lead to isomorphisms that map unrelated predicates that just happen to have identical assignments. However by additionally considering this possibility my result applies to other, less-constrained examples such as the canonical Suppes reduction model that is being targeted.

I see this as a problem because the case discussed seems to represent an excellent case for a successional reduction: theoreticians make a very naïve and simple improvement upon an existing theory. It seems obvious that a reduction from β to α is desirable, and, *ceteris paribus*, is hopefully an easy affair. Indeed, it is a criticism that is very similar to the one leveled against Kemeny and Oppenheim in §1.1.2, as it demands an amount of correlation between the two theories that is extremely stringent. A Suppes-reduction requires that α give *exactly* the same answers as β in regards to a given scenario that was in β 's domain.

It is an interesting question whether the Suppes archetype should be considered direct or indirect. An argument for its indirectness may be made on the grounds that the vocabularies of the theories may be different. But we still find that there must be an interface between the possible models of the theories, allowing the theories to still “talk to one another” – a hallmark of more direct models. For many who work in the sciences, an axiomatic representation of a theory is only interesting inasmuch as it generates models for the scientist to employ/examine. For those who work in general relativity, the models that satisfy the various field equations *are* the principle objects of study. Likewise for practicing engineers, the models of mechanics are what drive their results, not the baseline equations and force relations. Keeping cases like these in mind, it is tempting to consider the archetype to be a direct one. We have focused before on a connection between the theories' vocabularies when considering connections; a shift in focus to a connection between the theories' models seemingly should not change this classification significantly. Indeed, it can be shown without much difficulty that if β reduces to α by the Nagel model²³, that we will be assured of a Suppes-reduction of β by α (Schaffner 1967, 138).

This being said, recall that a potential problem for the Nagelian model brought up by Kuhn and Feyerabend concerned instances where some of the terms in each theory avoided translation due to incommensurability. Such a worry may be dodged here, as we are only concerned with the structural features of the models instantiated by the theories (Stegmüller 1976). These models must talk with one another, but such a conversation is held in a universal language that cares little about naming and all about form. So in a manner similar to what we saw for Kemeny and Oppenheim, a Suppes-reduction is indirect, in the respect that the theories don't really “talk to one another”, but instead merely compare pictures. On these grounds, I am

²³ To arrive at this proof, Schaffner needed to first formalize the conditions of Nagel's model. Although they differ slightly from the ones I have myself provided in §1.1.1, the same result will hold for my rendition of the Nagel model.

content to claim that the Suppes schema is in some ways direct and in some ways indirect. Searching for a more definitive classification is likely not possible; even if we were to have one I suggest that it would be rather unenlightening.

Despite the archetype focusing on the models entailed by the theories, we notice that the theories themselves minimally have to be the type that could be said to be satisfied by a reasonably-complex model. Much of the structuralist literature deals with theories that can be formulated set-theoretically in first- or second-order logic. What of the scientific theories that aren't overly mathematical, such as those found in ecology or biology? The theory of evolution gives explanations, but to imagine that we could give a semantic model that it would be satisfied by is perhaps a bit of a stretch. Furthermore, as a theory becomes more general and less precise, the number of models that will satisfy it explodes. So a theory of evolution, by glossing over the superfluous details in most cases, will result in there being a larger number of models that one is forced to address. In the abstract this does not appear to be much of an issue, but when attempting a reduction from evolutionary biology to a given α , the difficulties may arise that go beyond mere complexity. For instance, a lack of precision may make it difficult to judge whether certain fringe models in fact satisfy the evolutionary theory. And these may be precisely the theories that α cannot be satisfied by, extensions notwithstanding.

Notice that the archetype seems to be more applicable to successional reductions than inter-level reductions. This is just because the successor theory will usually be as successful as the succeeded theory, as this is typically why the newer theory was preferable to the old in the first place. However as we recently showed, certain types of improvements on an existing theory may in fact prevent **(I)** from being satisfied.

To summarize, the Suppes view of reduction focuses on the entailments of the theories involved, not the theories themselves. It is an indirect model that requires only a structural isomorphism of models: so long as every model that satisfied β could be, once extended, found isomorphic to a model that also satisfied α , Suppes's condition is satisfied. There need not be any relevant similarity of the vocabularies – or even subject matters – of the theories themselves. Lastly, considering that there may well be isomorphisms of models from disparate scientific theories, Suppes's requirement is often thought to provide only a necessary condition for a reduction.

§1.2 Goals and Aims of Philosophical Models of Reduction:

§1.1 gave us a detailed look at three classic models of reduction. To further assess these models, §1.2.1 will briefly consider two types of goals that could underlie a reduction: ontological and epistemic. §1.2.2 summarizes what each model allows. Each of the last two sections will permit a discussion about what aims each of the three authors had in making their models. By doing so, we will be able to see what general ideas underlie the three models, thereby motivating the subsequent accounts of reduction that began to emerge in their wake.

§1.2.1 Two Traditional Goals of Reduction:

To judge the adequacy of a reduction, be it a particular scientific example or a generalized account of an archetype, one must recognize what the *goals* of the endeavor are thought to be. Rarely does a researcher begin to compare two theories without some intended payout. Sarkar identifies two major types of goals that might be addressed by a reduction: *epistemological goals* and *ontological goals* (1992, 169). One of Sarkar's contributions was to highlight how it was these two interests might be conflated or confused, and that disagreements in the literature might be assuaged once the two types were understood as being (in most cases) distinct. A recognition of which goal is being pursued often serves to clarify a reduction's methodology, focus, and import. Sarkar speaks of two types of theoretical reductions that are driven by epistemic and ontological goals: *constitutive* and *explanatory*.

Sarkar defines a “constitutive reduction” as one that seeks to show something about the compositional structure of the scientific entities being examined. It hopes to exhibit a mereological character of the subjects involved, to tell how an often smaller, lower-level structuring of entities may account for the behavior of the entities at an often larger, higher-level structure. A hallmark of constitutive reductions is that they are essentially ontological in purpose: they seek to show how objects in the scientific landscape relate to one another, what the basic objects are, and how these objects may be said to compose the world.

Sarkar's "explanatory reduction" is, unsurprisingly, a reduction that has an explicit goal of explanation that involves "entities". Its goals are wholeheartedly epistemic²⁴ – it tells us the why/how/what of something in the world that needs explaining. An important necessary feature for scientific explanatory reduction, as I see it, is that the *explanans* be a scientific constituent. Perhaps, when reducing folk psychology to evolutionary psychology, we employ a facet of evolutionary psychology to provide an explanation for the otherwise unexplained folk-psychological fact that humans are inherently interested in art. Here we are showing a "given" of folk psychology to be a consequence of constraints imposed by an evolutionary theory of fitness. Regardless of the success of such a reduction, for it to be a reduction involving the sciences, science must take the role of explainer.

Theoretical reductions of all stripes may have goals that are neither ontological nor epistemic, however, as I hope to make clear in the chapters that follow. Sarkar notes that there are examples of theoretical reductions that exemplify one goal or the other, and that sometimes a theoretical reduction may have a purpose that is both ontological and epistemic (1992, 171). One of the major points made by Sarkar was to show how a recognition of which goals were operating helped clarify the discussion.

Sarkar was not the only writer to be concerned with goals; later we will see how William Wimsatt and others contribute. For the discussion that follows in §1.2.2, however, I will be content to rely on the two broad categories of "ontological" and "epistemic" goals. Even this simple distinction will allow a fruitful discussion concerning what the goals of a reduction may tell us about the aims of those who would seek to develop a model of reduction.

§1.2.2 Relata and Intention:

²⁴ Some philosophers distinguish between the terms "epistemological" and "epistemic". It seems clear that Sarkar intends the term "epistemological" to indicate a knowledge making or finding endeavor, as opposed to one that relates to a theoretical discussion about *how* knowledge might be achieved or attained. I will typically use "epistemic" to indicate the former relationship, and "epistemological" to indicate the latter, but it may be the case that I deviate from this convention because Sarkar and others do not see the distinction as relevant. On a related note, I will only make reference to the word "ontological" and not "ontic".

Each of the three models had “scientific theory” as *relata*. Granted Suppes focused on the models that satisfied a theory, but these are wholly parasitic on the formal theories which generate them. The possibility of reducing larger or smaller scientific elements was acknowledged, but each of the models was considered only theories. Also, for each of the models to be effective as I have extracted them, it must be that the theories are presented highly formalized; in fact, it seems that none of the models could get off the ground unless the theories themselves were first axiomatized (Schaffner 1967, 138). I see this as happening for several reasons. First, each account was introduced at a time when the literature concerning reduction had not significantly developed. Reduction of a theory seemed a natural starting point for any such investigation, in that it was one that could be cleanly explicated without need to worry about any complications added by “theory-parts” or “equations”. Secondly, I believe that the focus of each of the authors was in creating a model of reduction that maintained the idea that the authors each had for what a reduction should accomplish. In some respect this claim is trivial, but I believe it is more telling once we consider how strict each model actually is.

Kemeny and Oppenheim required that α explain *everything* that β did. This leaves open the possibility of a β that failed to explain some features adequately. But all that this β did explain, α is thereby accountable for an explanation. Take a given observation: it may be that β doesn’t adequately explain it, and α does. Here we have a clear preference for α . But suppose that instead β did explain our observation. Ignoring any difficulties in the comprehension or facilitation, we might just as well prefer α to provide the explanation. Thus we have a strong reason to want to dispense with β , or at the very least to recognize that, in essence, α is a better theory than β , in that it must have a larger explanatory range. If we are to understand “explanation” to mean “D-N explanation”, there is the further imposed requirement of having the equations of α and β agree *exactly* for all values.

Next examine Nagel. He was interested in showing how, once properly connected by bridge principles, we can have β be a consequence of α . This shows that α is capable of making *all* of the claims that β does, and more. Casting aside pragmatic worries of complexity, economy of vocabulary, or ease of methodology, this tells us that theoretically β is superfluous. β may serve an important historical role, and may be didactically helpful, but when concerned with the nature of the world, any preference for β over α is merely instrumental.

The Suppes conception of reduction similarly required that every model that satisfied β was isomorphic to a model (likely extended) that satisfied α . This again provides a picture of progress, as α has the ability to provide any of the “same answers” to questions that β could. α provides new predictions for otherwise unexamined domains; but where β speaks, α agrees exactly²⁵. This shows that science, as it progresses, merely increases its scope, for past theories may only be reduced if they were as correct as the successor theory was. This is a picture that has no room for error; excepting for mere luck in an isomorphism, we could never see the two structures coincide if they fundamentally disagree about a value.

I believe that each of these models of reduction supports a model of scientific progress²⁶, and do so very strongly. Each of the models had α “agree with” β for *every* in-principle relevant aspect; they differed merely in what they thought “agreement” should amount to. I have provided several criticisms of why I believe the requirements of the three models are too strong. Kemeny and Oppenheim require that α explain *all* that β did, which I believe is too stringent. Likewise, Nagel and Suppes required an *exact* agreement of values/predictions for every case. This also I feel is too restricting. Due to the high bars set by these models, the criticism would need to establish that many avowed paradigmatic reductions (as referenced by copious philosophers and scientists) would fail to qualify as “reductions” under the three models we have focused on in this chapter.

My purpose in this section has been to set the stage for this criticism to be received. Indeed I will not be the one to provide it; other philosophers will notice the deficiency of these models and attempt to either revise them (§2.1), or engage with theory reduction/comparison in a new way (§3.1, §4.1). The next three chapters will continue this narrative in relation to specific scientific case studies.

§1.3 Conclusion:

²⁵ If we recall the theorem proved by Shaffner that a Nagel reduction implies a Suppes reduction, this claim is obvious.

²⁶ This idea of scientific progress has sometimes been conjoined with a realist worldview, oft referred to as “convergent realism” (Laudan 1981). We will elaborate on this position in §6.

The purpose of §1.1 was to introduce three major conceptual models of reduction. Nagel's direct model required a derivation of the reduced theory from the reducing. Kemeny and Oppenheim championed an indirect model, allowing for a reduction if the reducing theory had the ability to better explain known observations than the reduced. Lastly Suppes maintained that the models that satisfied theories be the focus of a reduction: it was a necessary condition for a reduction that any model that satisfied the reducing theory must also satisfy the reduced. §1.2.1 mentioned the importance of goals when regarding these models. A reduction may be constitutive, relating to the dependencies of scientific entities, or may be explanatory, occurring when the reduced theory is explained by the reducing theory. These two types of reduction relate to ontological and epistemic goals. It is additionally fruitful to postulate about the intentions of those who made the models of §1.1. §1.2.2 shows how each of the three models tacitly rely on an idea of scientific progress that is unable to account for error. Much of the criticism levied against the three views in §1.1 hinged on this inability.

As a whole, this chapter presents early views on reduction. It showcases how philosophers first began to consider intertheory activity, and what they believed to be important characteristics of successful reductions. As the next few chapters will show, progress beyond these initial models issues from a desire to accommodate approximation and lack of exactitude.

Chapter 2 - Analogous Theories: General

Relativity to Newtonian Mechanics

What has been emerging from re-studying historical examples is some appreciation of the almost bewildering variety and complexity of what are regarded as reductions.

Clifford Hooker (1981a, 40)

§2.0 Introduction:

In the wake of the shortcomings of the reduction models of §1.1, Schaffner developed his own model of reduction to accommodate cases where the reduced theory could be false. §2.1 will detail the Schaffner model: one where the reducing theory derives a *modified* reduced theory that will be analogous to the reduced theory proper. Subsequent work by Clifford Hooker, Patrica Churchland, Paul Churchland, and John Bickle constitute the “New Wave” of reduction analysis developed in recent years. As §2.1.1 will show, the New Wave alters the model provided by Schaffner in several ways, notably by requiring that the analog theory be constructed from the conceptual resources of the reducing theory, not the reduced.

§2.2 describes the limiting processes used when relating general relativity (GR) to Newtonian mechanics (NM), originally discovered by Andrzej Trautman. We will try to fit this case study to the models of §2.1, showing how well a contemporary successional reduction will fit such models. The other function of the case will be to highlight what goals have been realized by the reduction. §2.2.1 shows traditional epistemic goals of one theory explaining another, while §2.2.2 gives evidence of another type of epistemic goal, where the successor theory explains the historical success of the succeeded theory.

§2.1 The Schaffner Model:

Nagel's model of reduction from §1.1.1 – indeed each of the three models of §1.1 – demands that there be an exact correlation between the predictions of theory α and theory β . This is problematic because it disallows any error in β . Examine a prediction of general relativity: an observer on the earth's surface will observe that clocks measured from an airplane above the earth's surface tick slightly faster than clocks on the surface. As Newtonian mechanics predicts that there will be no such time discrepancy, there can never be a Nagelian reduction between the two. This is because Newtonian mechanics is *false*, as verified by observations. Derivation preserves truth while bridge principles are empirically-verified biconditionals; inasmuch as we are to assume that GR is true, it is thereby entailed that GR's deductive consequences must likewise be true. For an intra-level reduction to be possible under the Nagel model, it would have to be that the successor theory α had a larger scope than the succeeding theory β , as error cannot occur in β .

Rarely do we see instances of intra-level reductions where one theory exactly correlates with its successor, and the two differ merely in the range of phenomena they seek to cover (Sklar 1967, p.110). Instead we see theories overturned by other theories that often have significantly different worldviews, successor theories that recognize the falsity of their predecessors. This is why we talk about “scientific revolutions” and not “scientific extensions”. A similar situation occurs for inter-level reductions. When the “bigger” entities of β are composed of the “smaller” entities of α , often there are discrepancies when we talk about the properties of the wholes *qua* wholes and the wholes *qua* sum-of-parts. For both intra- and inter-level reductions, we generally find that α *corrects* β .

To fix the problem posed by the demand of exact correlation, and the related worry of limited application to real-world cases, Schaffner makes significant revisions to the Nagel model²⁷. Schaffner proposes a model that is still derivational, yet the derivation does not occur

²⁷ For some the value of the Nagel model need not lie in its applicability to actual scientific cases. Nickles claims that he “would go so far to say that Nagel's derivational reduction is a useful concept even if not one single historical example perfectly exemplifies the pattern he describes” (Schaffner 1974, 185).

directly between α and β . Instead we use a theory which is very similar to β , called β^* . Here I provide what I believe is an accurate representation of Schaffner's model²⁸:

β is *reduced* to α iff:

- (I) The individuals/predicates of β^* are linked with individuals/predicates or groups of individuals/predicates of α by empirically-supported reduction functions.
- (II) β^* can be derived from α by means of these reduction functions.
- (III) β^* corrects β , in the sense that:
 - (i) β^* provides more accurate predictions than β in most cases.
 - (ii) β^* explains why β was incorrect, and why it was once considered correct.
- (IV) β^* and β are linked by a relation of strong analogy, $[A_s]$.

As with Nagel, we have bridge laws (at (I)) and derivation (with clause (II)) facilitated by these bridge laws. However there is a notable difference in what is derived: previously the goal was to derive the reduced theory β from α via some bridge laws. Now we are to take α along with bridge laws to derive β 's analog, β^* . β^* is thought to be like β , in that it is a modification of β that still possesses key concepts of the original. The relationship, $[A_s]$, between β and β^* is less-clearly defined. Schaffner explains that β^* "bears a close similarity to β and produces numerical predictions which are 'very close' to β 's" (1967, 145). The imprecision of this relation can be considered one of the model's strengths, as it does not put a strict requirement on how closely related β and β^* must be to one another. When $[A_s]$ is the identity relation, the reduction process is exactly the Nagel model of §1.1.1 (Schaffner 1967, 145). In this case, any mention of β^* is superfluous, for a simple derivation from α to β would suffice. Thus we may view Nagel's model as a special case of Schaffner's.

Although the details of how the reductive functions of (I) play out are not made explicit in my formulation, Schaffner provides the following summary:

All primitive terms of β^* are associated with one or more of the terms of α such that:

²⁸ To arrive at this extraction, I have relied on several of Schaffner's accounts (1967) and (1974). I have mostly maintained the ordering of (Schaffner 1967).

- (a) β^* (entities) = function [α (entities)].
 (b) β^* (predicates) = function [α (predicates)]. (1974, 618) [2.1]

Recent work by Rasmus Grønfeldt Winther represents the reductive functions by the following equation:

$$\text{Term}_{i,\beta^*} = \text{function}_i (\text{Term}_{i,\alpha}), \text{ where } 1 \leq i \leq n \quad (\text{Winther 2009, 124}) \quad [2.2]$$

Winther also claims that the reductive functions of the model are bijective, revealing why he required there be n terms in the domain and range²⁹. This seemingly contradicts what Schaffner has claimed in [2.1] in which requiring “terms of β^* are associated with one or more of the terms of α ”, as this could not happen were the function to be one-to-one, not many-to-one. To resolve this difficulty we will have to be very careful as to the meaning of “terms”.

Winther intends the “terms” of [2.2] to be composed “of entities, predicates, or relations, or combinations thereof” (2009, 124). We then infer that a term may be singular or composite. For example, assume that α is “Standard Model particle physics” and β^* is “corrected atomic physics”. We can imagine a reductive function as follows³⁰: $F(\text{up-quark, up-quark, down-quark}) = \text{proton}$. Here we have a many-place function of a group of individuals, and one that seems necessary to translate between the two theories. A many-placed function that links terms from α to terms of β^* and is bijective must have the same number of terms in the domain and range. Yet let us be clear that this need not require an equicardinality of entities in α and β^* . This is because the domain is here populated by three-tuples, with each three-tuple composed of individuals from α . “Many-to-one” is thus slightly ambiguous: it might represent either a single-place function that is not injective, or a many-place function that is either injective or not. Winther (and Schaffner) intend the latter: a many-place function that is injective. This is why [2.2] says “term” and not “terms”. For Schaffner’s [2.1], I think that “term” should be taken to mean “individual” or “predicate”, making it equivalent to the formulation of (I). This implies that Schaffner

²⁹ Although a small point, I believe that Winther should have made the subscript i of “Term _{i,α} ” instead a j . Even if the number of terms of α and β is equal, the way which we have numbered the terms β need not correspond to the way that we have numbered the terms of α .

³⁰ Another option would be to have a single-valued function whose range consisted of logical compositions of entities of α , making the picture instead be: $F(\text{up-quark} \ \& \ \text{up-quark} \ \& \ \text{down-quark}) = \text{proton}$. Changing the representation to look like this would only trivially alter the discussion that follows.

believes that the reduction function's domains are populated by individuals or groups of individuals, whereas the range must be solely individuals³¹. In this respect it is consistent with Winther's extraction, if not a bit more restricting.

As if in response to the models examined in §1.1, (III) makes two novel demands. First is the stipulation that there be a quantitative improvement of β^* over β in *most* cases. I believe that this is preferable to demanding that β^* in *all* cases offer better empirical results than β . As the criticisms of §1.1 urged, requiring universal improvements eliminates many cases that otherwise might well qualify as a reduction. Verifying that (III)(i) obtains is an experimental endeavor. If we are attempting to apply the Schaffner criterion to a successional reduction, we will likely already assume that (III)(i) holds, as it is presumably on the basis of empirical merits that α is deemed a “successor” to β .

(III)(ii) is an explanatory requirement, one that did not receive much focus in subsequent discussions of the model by critics (Hooker 1981a) (R. P. Endicott 1998) (Winther 2009), or even later in the examples which Schaffner himself provides. I find this the most interesting clause of the model, as it explicitly gives some hint at what goals there could be for a reduction. Later, when discussing a restatement of his model in 1974, Schaffner again includes the requirement that “the reducing theory explain why the uncorrected reduced theory worked as well as it did historically, e.g., by pointing out that the theories lead to approximately equal laws, and that the experimental results will agree very closely except in special extreme circumstances” (1974, 617). This extremely pithy summary receives no further attention in his article. This is regrettable, for the short statements provided by Schaffner are in need of some elaboration. The characterization of “approximately” and “very closely” are left vague, a fact of which Schaffner himself is well aware. Also, how does one decide what makes a given circumstance “special” and “extreme”, and what makes another circumstance germane and unfit to qualify as an exception?

Other authors have echoed the requirement of (III)(ii). Lawrence Sklar distinguished between two types of reduction: cases where α *explains* β , and cases where α *explains away* β (1967, 112). The former is the traditional epistemic goal highlighted by Sarkar in §1.2.1, where α explained features or constituents of β . The latter case differs in its *explanandum*: here we find

³¹ The usage of examples in Schaffner's response to Hull concerning the stipulation of reduction functions for complex biological notions such as “dominant” and “epistatic” further corroborate this intention (1974, 621–622).

that α explains β 's "apparent success" (Sklar 1967, 112). Presumably thinking of some correspondence-principle-inspired reduction, Sklar provides the example of quantum mechanics explaining not "Newtonian mechanics", but: "why [Newtonian mechanics] met with such apparent success for such a long period of time and under such thorough experimental scrutiny" (1967, 112). Unfortunately, Sklar too stops short of telling us how such an explanation would happen in detail. Wimsatt, citing (Sklar 1967, 112), summarizes this as a case where α "explains why we were tempted to believe" β (Wimsatt 1974, 687). The actual mechanisms that underlie this second type of explanation have not been much discussed, and no detailed examples have been given. In direct reaction to this deficiency, I will meet these issues head on in §2.2.2, and additionally in §3.2.1, §3.3.1, and §3.4.1.

§2.1.1 The New Wave Model:

The picture of reduction offered by Schaffner's model has been modified in recent years in the work of Hooker, Churchland, Churchland, Bickle, and Wimsatt. Together, the prevailing model has been referred to as the "New Wave"³² (R. P. Endicott 1998) (Eck, Jong, and Schouten 2006). Their work has developed many of the ideas originally under-specified in Schaffner's model, and has also modified others. The first significant change involved the nature of the analog theory. Previously, β^* was a modification that was made to β . It was "close" to β , yet it was grounded in the theoretical resources of β . This is why there needed to be bridge principles that linked α and β^* , because the derivation was inter-theoretic. Taking their cue from many scientific examples, the New Wave saw that it was instead more appropriate to make an approximate theory out of the conceptual resources provided by α . The analog theory simulates β on α 's terms, and is to be "expressed in the vocabulary proper to α "; it is the "equipotent image of β within the idiom of α " (P. M. Churchland 1985, 10). Granted this different role, and to help distinguish the New Wave model from Schaffner's, I will refer to the analog theory of the New Wave as α^* .

³² There is naturally a worry that there are differences in how each author's account proceeds, as some have made explicit (Endicott 2001). I will avoid a detailed exegesis of the different positions and instead attempt to summarize what I believe is their common core.

Previously the Schaffner model took α and attempted to derive β^* , with some help from the bridging reductive functions. The New Wave, looking to have α entail α^* , also seeks aid in the form of a specific set of circumstances, C . These circumstances are thought to be the particulars that one would instantiate within α , or more generally factors that acted upon α . Certain subcases of α fall under the purview of C : the times when we want to restrict the theory to cover cases of “only perfect energy conservation” or “perfect competition without arbitrage”. C may consist in a list of biological minutiae of a “fearful complexity” (Hooker 1981a, 49), could be counterfactual, or could instead be a set of constraint parameters that operate on the variables contained in α . Whatever additional factors are introduced at the onset of the procedure, aside from the original, are represented by C .

By recognizing C , the New Wave acknowledges the fact that most reductions occur with regards to specific scenarios or modeled contexts. Statistical mechanics could not operate without its boxes and walls, for example. Here any comparison between heat and kinetic energy of molecules is to be conducted in regards to specific modeled paradigms. A number of spherical molecules are chosen and then placed under inelastic volume constraints – these are the “walls”. So long as the ball population is sufficiently large, the resultant relationships may be confidently manipulated and limited so as to arrive at macroscopic heat equations³³. It is important to note that C is intended to either manipulate or constrain α . Thus worries about a contradiction with α , thereby entailing any proposition, are minimal.

Once C has been imposed on α , we arrive at α^* , which is again to be related to β by analogy. This analog relation, which I will call $[A_n]$, compares α^* and β inter-theoretically. This is different from the $[A_s]$ of Schaffner, which was an intra-theoretic analogy. $[A_s]$ and $[A_n]$ are intended to fill a similar role in each model. In Schaffner’s model, for the derivation to follow there had to be bridge laws to connect the disparate vocabularies (illustrated by **(I)**). In the New Wave model, the manipulation of theory concerns the transformation of α into α^* , making the connections confined to the theoretical vocabulary of α . As all inter-theoretic comparisons of the New Wave account are subsumed into $[A_n]$, there need not be any explicit bridges between α and β . Instead we merely require that α^* “acts like” β does, in some fashion. Thus the New Wave

³³ As has been the gradual trend, reductive stock examples are fleeting. The actual effectiveness of various reductive strategies in thermodynamics has been criticized by Sklar (1995).

dispenses with a need for bridge laws, or any of their associated worries (including those mentioned in §1.1.1 for instance).

One way of representing the New Wave model comes from Hooker (1981a, 50):

$$\beta \text{ reduces to } \alpha \text{ iff: } (\alpha \ \& \ C) \rightarrow \alpha^*, \text{ and } \alpha^*[A_n]\beta \quad [2.3]$$

In some cases, I think that viewing C as a group of propositions that are simply conjoined to α is correct. In a cosmological reduction, adding a C of “space is infinite” is merely telling us one more condition that will factor into the deduction of α^* . However there are other processes that occur in reductions that are not as easily captured by mere conjunctions.

Hooker refers only to cases where α and α^* are “close” and makes little reference to any particular way that this notion of closeness may be made more precise. Churchland and others recognize that the modification of α may involve *limiting* (P. M. Churchland 1985, 10), one way of mathematically capturing how two functions may be said to be “close” to one another. This is an important process for many reductions: §2.1.2, §3, and §4 will involve a detailed discussion of how limiting is involved in inter-theory relations. Assume that a term³⁴ ϵ is employed significantly within α . It is far from obvious to see how the limiting of ϵ can be represented by conjoining additional propositions about ϵ to α , a move that is made by Churchland (1985, 10). Mathematical operations such as addition and subtraction are well-known to have a representation in purely logical terms. However little work has been done to see if a claim such as “ ϵ limits to ϵ_0 ” can be captured by first-order logic. To assuage this worry, it might be prudent to instead view the New Wave model of reduction as claiming:

$$\beta \text{ reduces to } \alpha \text{ iff: } \alpha^C \rightarrow \alpha^*, \text{ and } \alpha^*[A_n]\beta \quad [2.4]$$

Here we are to read α^C as: “ α with the conditions C imposed upon it”. I mention this only as a precautionary measure: if we indeed can represent processes such as limiting of an equation by a conjunction of propositions to that equation, then I am content to concede that [2.3] is an adequate representation and [2.4] is unnecessary.

³⁴ I have avoided calling ϵ a “variable” or a “constant”, as such a distinction is controversial and important, as §3.1 will demonstrate.

As before, there are ways that the New Wave reduction process is very similar to the one espoused by Nagel. Let $\mathbf{C}$ become the empty set, requiring that α^* be equal to α . Now all of the weight of the reduction falls on $[\mathbf{A}_n]$. This might be too much to bear, for it merely turns reduction into a case of analogy. And as Nagel's model required that derivation link α and β , so long as we grant that entailment is sufficient for an $[\mathbf{A}_n]$ the two models will function identically. But there are many cases where the difference between the two theories falls somewhat short of exact deductive entailment. The question of how close α^* can be to β asks for some philosophical characterization. In his 1981 article, Hooker notes that "satisfactory criteria for an $[\mathbf{A}_n]$ warranting reduction have never been offered, indeed rarely broached" (1981a, 50). Hooker then briefly provides two factors himself: (1) normative commitments and (2) experimental capabilities. I find both very intriguing suggestions, and will examine each in some detail. When discussing the first criterion, Hooker claims:

The meta-theoretical, normative commitments of scientists are involved in assessing the adequacy of an $[\mathbf{A}_n]$ to warrant a particular reduction. Consider, e.g., debates about which features of theories, or of theoretical models of the world, most need to be preserved in reduction... The reasonableness of the $[\mathbf{A}_n]$ adopted must be argued in the light of the meta-philosophy stated. (Hooker 1981a, 50)

Decisions about what is theoretically important temper scientists' judgment of whether two theories are related closely enough. Perhaps a β describes temperature as being continuous. Assume that, after deriving α^* from some α , we are left with an α^* that has temperature discrete or quantized. Can we say that the two are enough alike to warrant a reduction? How flexible is the $[\mathbf{A}_n]$ in this respect? Hooker claims that this is a complicated judgement that will have to be made by the practitioners immersed in the study of heat. Whatever biases they have, including ones that are indeed meta-theoretic³⁵, become pertinent.

The second criterion concerns the "state of technological development and practice" (Hooker 1981a, 50). This would seem to be a reference to how accurate our empirical resources are when assessing or measuring constituents of a theory. Were we to find that string theory

³⁵ Although a small point, I disagree with the locution "metaphilosophical". I think that the scientists making decisions about two theories tacitly do so on commitments to what is scientifically valuable, but I doubt in many cases that these are philosophical commitments. In some instances, perhaps even in the discussed nature of temperature, I am perhaps willing to admit that they are indeed philosophical, as the discussion of temperature could depend upon how one wants to conceive of the nature of the phenomenon of "heat" (Chang 2007). However I do not see how any of this could be construed as "metaphilosophical".

(taking this as our α) predicted a mass for a particle that the standard model (this is the β) employed with slightly different mass, considering our current experimental capabilities we may be justified in making an analog of the standard model on string theoretical terms (α^*). This analog would have the reduced target mass predicted by α and we would then be able to capture it in a reduction. The relation $\alpha^*[A_n]\beta$ would be satisfied because experimentally we are uncertain about the true value of the mass of the particle.

The New Wave advocates have focused less on how to decide whether $[A_n]$ obtains, and more on what follows once it has been decided. The move here is to let the strength of the analogy dictate what type of ontological/theoretical relationships obtain between α and β . Say that there was a strong analogy between α^* and β , one where there was little difference in the terminology/configuration of the two theories. The confidence in this connection then represents a *retention* of ontology and major theoretical edifices. On the other end, were there to be a loose $[A_n]$, the thought is that there would be a disagreement about the mechanisms and entities of α^* and β as they functioned in the description of the phenomena. Loose analogies likely cover situations where the reduced theory is “radically false”, and little of the relevant constituents of β correspond with those of α^* . In these cases, we give preference to the ontological and theoretical framings of the reducing theory and eliminate the ontology or conceptual moves of the reduced theory; here we are left with an instance of *replacement*.

As there is a spectrum of strong and weak analogies that can potentially occur between the various α^* s and β s, there is likewise a spectrum of ontological commitments that follow from reduction, creating a *retention-replacement continuum*^{36, 37}. Furthermore the strength of an analogy is also likely to provide an indication of the epistemological ground covered from α to β : the more of a stretch the analogy is making, the more our knowledge about the world has shifted. As opposed to spectrum of “strong” or “weak” analogies, the New Wave will at times refer to analogies that map β to α^* in a manner which varies from “smooth” to “bumpy” (P. M. Churchland 1985) (Wimsatt 1974) (P. S. Churchland 1989). John Bickle provides an excellent summary of the position: “We can lay out this spectrum (informally) as follows: at the left-most

³⁶ In 1977, Schaffner expanded his model so as to create a similar spectrum, albeit one whose range was dependent upon $[A_s]$ rather than an $[A_n]$ (1977). In §2.1.1 I have omitted to cover this progression and instead focused on the Schaffner model in its original formulation.

³⁷ Endicott argues that attention to type reductions and token reductions will lead us to amend this spectrum (Endicott 2007). I have not included a discussion of this amendment as I believe it doesn’t significantly contribute to my discussion in this chapter.

endpoint lie the perfectly smooth reduction pairs, where α^* is the exactly equipotent isomorphic image of β ...at the right-most endpoint lie the extremely bumpy or replacement cases. And separating these two extremes is a continuous spectrum of cases, approximating more or less closely one of the two endpoints” (1992, 417).

The New Wave contributors recognize that theories are rarely cold and unresponsive entities (Hooker 1981a) (Wimsatt 2007). Instead, they are dynamic reactive components of active scientific practice. As such, the reductive process may be actively involved in the development of each theory. Often theories progress in tandem: by reactively modifying each other in relation to reductive goals between the two, there is a theoretical “co-evolution”. Here there is a feedback generated by both the non-trivial act of creating α^* from α with C , and the complicated process of considering the subsequent $[A_n]$. β may evolve from this comparison into a new theory β' . Now, when a new reduction is attempted, something must change. Sometimes it is the addition of a new C' to α to create a new $\alpha^{*'}.$ In this case the changes that result from postulating the necessary $\alpha^{*'}.$ can result in us making subsequent changes to α , creating a new theory, α' . Other times C will remain relatively the same and we must directly postulate a new α' which can result in a modification to α . Finally, we may have to modify both α and C . Here each case results in a modification to the reducing theory, creating an α' . This evolution may continue, leading to feedback for β' which results in the formulation of a β'' etc. The existence of this process in the sciences shows how a reduction need not be epiphenomenal to the act of scientific theory-making. Reductions are certainly important for demonstrating ontological and epistemological considerations for the philosophically minded to engage with; however a co-evolution of reducing theories would demonstrate that reduction is an important tool with which to *improve* theories, as opposed to a mere aid in understanding or characterizing them.

As we have seen, there are many similarities between the New Wave model and the Schaffner model. To recap, I will repeat the Schaffner model’s criterion:

β is *reduced* to α iff:

- (I) The individuals/predicates of β^* are linked with individuals/predicates or groups of individuals/predicates of α by empirically-supported reduction functions.
- (II) β^* can be derived from α by means of these reduction functions.
- (III) β^* corrects β , in the sense that:

- (i) β^* provides more accurate predictions than β in most cases.
- (ii) β^* explains why β was incorrect, and why it was once considered correct.
- (IV) β^* and β are linked by a relation of strong analogy, $[A_s]$.

Both models may account for cases where one theory is approximate to another, both use deduction from α to a corrected theory, and both involve analogy between a corrected theory and β . They differ in what they think the corrected theory should be constructed from (namely, the conceptual resources of α or β), the necessity of bridge principles, and whether the analogy is inter-theoretic or intra-theoretic. Looking at Schaffner's model, this serves to compare each of his criteria, *with the exception of (III)(ii)*. The little-discussed explanatory requirement is notably absent in the New Wave model, an absence which §2.2.2 will elaborate upon.

§2.2 Trautman's Reduction:

Now I will provide an extended example that fits the picture of reduction championed by the New Wave. By providing this example, my hope is to elucidate how a philosophical model of reduction may be applied to a case from contemporary science, as well as to bring attention to the roles played by its particulars. The other payoff of such an example is that it will provide the details of how a reduction may serve several goals, as we will see in §2.2.1 and §2.2.2. A strength of the New Wave model lies in its ability to apply to situations that seem quite complicated. Much of the weight of a successful reduction, as will be shown, is shouldered by $[A_n]$.

We take our cue to examine such a case from Hooker. He mentions a case of a successional reduction from physics that, by his consideration, fits the New Wave model. I will reproduce his discussion in full:

For example, there is a α^* for local, macroscopic space and time derivable from General Relativity Theory (GR) which is closely analogous to Newtonian mechanics (NM); objects (states etc.) in NM can be identified with objects (states etc.) in GR provided that, in addition to the localness assumptions characterizing α^* in this case, additional constraints on accuracy of measurements, energy densities etc., are also imposed. The

informal tendency has been to identify α^* and β and ignore these complications, thereby introducing much confusion. (Hooker 1981b, 203)

The case Hooker refers to was originally presented by Trautman (1965) and later given a contemporary presentation by David Malament (1986a). In the 1986 article, and in his subsequent book (2012), Malament is reluctant to deem the case a “reduction”. James Owen Weatherall is likewise cautious to officially weigh in on the debate when he discusses the result (2011). It is an open question as to whether we consider the work that follows a reduction; for the remainder of the section I will follow Hooker and refer to it as one³⁸. Also, it may well be that Hooker did not in any way have the Trautman result in mind in the above quotation, as he believes that to arrive at α^* we needed to add “additional constraints on accuracy of measurements”, something that we shall see will not factor in to our considerations for creating α^* . If this were so, I believe it makes the work in this section even more so an interesting application to the New Wave model.

This case is compelling because the discussed α^*/β decomposition does much to clarify the details of the reduction’s structure. Allow me to briefly describe the initial stages of how the reduction proceeds. First we take a class of GR spacetimes as our α , represented by a manifold, $\mathbf{M}$, and a metric, $\mathbf{g}_{ab}$. Each manifold is four-dimensional, smooth, and everywhere connected. The metric allows us to assign lengths to vectors. Importantly, the metric will also allow us to describe how light-cones open at any given point. Furthermore, the $\mathbf{g}_{ab}$ is directly related to the distribution of mass in the universe; Einstein’s equation explicitly provides the details of how the mass-density affects spacetime curvature. Note that not any metric will do: each $\mathbf{g}_{ab}$ satisfies Einstein’s equation. Indeed, this is why it is a theory of “gravitation”, for objects “influence” the motion of other objects, in that they dictate how space unfolds over time. This α , formally the set of all GR spacetimes, $(\mathbf{M}, \mathbf{g}_{ab})$, will be parameterized about a variable, λ . λ , as embedded in $\mathbf{g}_{ab}$, dictates the maximum speed which a particle may travel, an important concept for GR. Decreasing λ will increase the maximum particle speed, causing the light cones at each point to “spread”. We take the limit $\lambda \rightarrow 0$ on $(\mathbf{M}, \mathbf{g}_{ab}(\lambda))$; this is the process that we will consider to be the constraint, $\mathbf{C}$. As the parameter approaches zero, the light cones will open wider and wider (effectively allowing the “speed of light” to be unbounded). The end product of limiting $(\mathbf{M},$

³⁸ The issue of whether we should consider the comparison done by Trautman a “reduction” will again receive some discussion in §4.1, when Batterman seeks to distinguish between the types of limit employed in this case and others.

$\mathbf{g}_{ab}(\lambda)$) is a class of classical spacetimes that, by virtue of possessing certain curvature conditions, are referred to as “Newtonian” (Malament 1986a). Specifically, as the light cones have been flattened, there is no upper bound limiting the speed of objects. Newtonian classical spacetimes are geometrical, as the motion of objects is a product of the geometry of the spacetime, just as it is in GR. These classical spacetimes are our α^* . And clearly we would like β to be canonical NM: a non-geometrical theory of dynamics whose equations of motion dictate the influence on objects by a gravitational potential field in Euclidean space. Naturally, it turns out that β and α^* are related, but the $[\mathbf{A}_n]$ that is tacitly expressed is rather complicated. All the components featured in [2.3] (or in [2.4]) are thus introduced.

Notice how this case cannot easily be said to fit the Schaffner model. Under the account provided in §2.1, we required that the analog theory, here called β^* , be constructed out of the conceptual resources of β . For this case, doing so would require that we somehow modified NM to make a theory that could be said to be analogous to α . It is difficult to see how this could happen, as there is significantly more structure provided by the conceptually more-sophisticated GR. For instance, take the basic fact of simultaneity being relative to reference frame, a property that GR spacetimes possess. I cannot see how this concept, and the quantitative observational predictions that accompany it, could be replicated within a transformed NM. Any such move would simply have to tack on very unintuitive stipulations about how an accelerated reference frame would behave. If this could even be done cleanly it would seem grossly *ad hoc*. Furthermore the $[\mathbf{A}_n]$ linking this hypothetical β^* and β wouldn’t be strong, precisely because this new β^* has a significant component that doesn’t really look anything like β . Thus any formulation of a “geometrized NM” would have to proceed from the conceptual vantage-point of GR, not NM.

There is more to be said about the classical spacetimes that comprise our α^* : although the geometrized version of Newtonian gravitation was originally formulated in the 1920’s by Elie Cartan³⁹, it should be emphasized that it could just as easily have been conceived of by Newton in the 1700’s had he access to some more advanced topology and differential geometry. That the geometrized Newtonian theory was developed in the wake of GR is historically relevant, yet not in any way essential – the two theories are conceptually independent. Likewise, much of the work done in developing the geometrized version was for the purpose of gaining a better

³⁹ Malament provides a nice list of references that concern the historical details (1986a).

understanding of how GR and NM relate. However since that time the geometrized Newtonian theory has been employed in answering extant difficulties that were otherwise quite vexing to answer in the standard framework⁴⁰.

Historical interests aside, how the mathematical relationship between β and α^* unfolds is much more subtle. Each theory provides a description of how objects will move, yet we are interested in when the descriptions can be made to fit one another exactly, that is, when the motions are identical. Given a non-geometrical formulation of Newtonian gravitation, it turns out that we are able to find a geometrized formulation that will describe the unfolding of events identically (Malament 1986a, 191). What is more interesting for the reduction at hand is that it is possible to go in the other direction: given a geometrized version of Newtonian gravitation it is possible to create a non-geometrized version such that each accounts for an event-structure accurately.

The result which allows us to “recover” NM from a geometrized theory is the *Trautman Recovery Theorem*, fully explained and proved in technical detail by Malament (2012, 4.4.5). For our purposes, it will do to notice that only certain geometrized spacetimes can be mapped to NM. As these are a portion of α^* , let us call this group $\alpha^{*'}.$ Not all spacetimes from α^* can be input for the recovery; instead only specific ones from $\alpha^{*'}.$ that possess the right curvature conditions will work⁴¹. These mapping subtleties provide us with the first interesting mathematical feature to notice about $[A_n]$.

Secondly and perhaps more noteworthy is that, when performing the recovery from $\alpha^{*'}.$ to β , the mapping is no longer unique. Given a classical spacetime of the class delimited by $\alpha^{*'}.$, we will be able to find an analogous Newtonian framework that, via a gravitational vector field that interacts with the mass distribution, possesses the same dynamics. However it happens to be that there are many other such Newtonian frameworks that are likewise identical. The reason for this is rather straightforward. Under our $\alpha^{*'}.$ spacetime, the motion of objects is geometrically prescribed, yet when employing a β framework the motion instead must be accounted for by a

⁴⁰ Here I am thinking of the success of the geometrized version of Newtonian mechanics in addressing issues concerning infinite spaces filled with infinite particles (Malament 1995) (Malament 2012, chap. 4.4).

⁴¹ The resultant classical spacetime must be either Friedmann-like or asymptotically flat. I direct the reader to (Malament 2012, 444). Specifically, the recovery theorem can only apply if $R^{ab}_{cd} = 0$. We know that the limited group has $R^{abcd} = 0$, spatial flatness. Yet some of these spacetimes have $R^{ab}_{cd} = 0$ and some have $R^{ab}_{cd} \neq 0$. Thus, only some of the limited classical spacetimes may be recovered. There is a class of spacetimes that are in the limited group that cannot be shown to be equivalent to a non-geometric spacetime, namely, those that have $R^{ab}_{cd} \neq 0$.

gravitational potential, and furthermore this must interface with a derivative operator. By doing so, it turns out that there are an infinite number of distinct β -type frameworks that may be recovered that each describe motion just as the single given classical spacetime of the sort dictated by $\alpha^{*'} did. Simply put, we lose uniqueness when formulating a NM-partner for each classical spacetime⁴².$

{Figure 2.1} summarizes the components and details of the reduction, with the following assignments:

α : All General Relativistic Spacetimes (M, g_{ab})

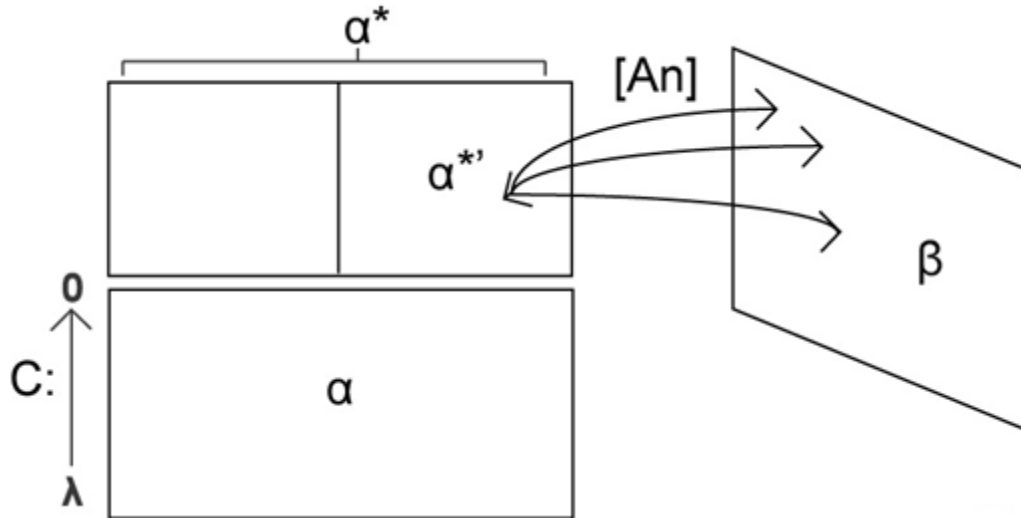
C: The Limit $\lambda \rightarrow 0$ on $(M, g_{ab}(\lambda))$

α^* : Class of Geometrized Newtonian Spacetimes

$\alpha^{*'}$: Specially-Curved Geometrized Newtonian Spacetimes

$[A_n]$: One-to-Many Mapping (Trautman Recovery Theorem)

β : Non-Geometrized Newtonian Mechanics



{Figure 2.1}

Each mathematical caveat has interesting consequences. In the first case, we have shown that the reduction is not encompassing. Only certain types of classical spacetimes from α^* have partners in β – precisely those that are in $\alpha^{*'}$. Specifically, they are ones that involve a certain

⁴² I refer the reader to Malament for these details (2012, 4.4.5).

amount of symmetry or behave very uniformly. The class of GR spacetimes that will limit to such classical spacetimes and yet cannot be further reduced to NM (i.e. those that are in α^* yet not in $\alpha^{*'}$) is therefore not in any way an aberrant one, making the excluded class significant and thus not easily dispensable. The second caveat shows yet another limitation of the reduction: as we lack uniqueness in our mapping we can no longer describe which specific formulation of NM is recovered. Therefore it stands as a worry that we will have inextricably “recovered” too much, as such a recovery is overdetermined: there are an infinite number of possible NM-arrays with differing gravitational potentials occupying β and each will adequately account for the motion of objects as described in α^* . Regardless of how such mapping details tax the relation, these details supply important facets of our $[A_n]$. One function of $[A_n]$ is to show how β and $\alpha^{*'}$ are structurally similar. In this specific case it should be clear that a mere structural mapping is insufficient.

The above mathematical intricacies are sure to be the focus of any physics journal that discusses them. However, I suspect that to the philosophically-inclined reader they will take second shelf to the interesting conceptual differences between the two theories. Newton conceived of a gravitational force pushing and pulling between objects so long as they possess mass: mass causes objects to move amidst a static background of space, and nothing else. For Newton mass did not “deform space” or anything to this effect. Furthermore, space is Euclidean and time is a variable independent of space. In contrast, in the classical spacetimes picked out by $\alpha^{*'}$, objects have their motions described by how the mass field interacts with space and time as a conglomerate. There is no “push” or “pull”, just motion along smooth gradients. Events in spacetime “unfold” relative to one’s worldline, yet their occurrence was already dictated by the metric structure embedded on the manifold. Indeed, even the notion of “spacetime” is quite distant from Newtonian physics *qua* eighteenth-century-physics practice: “the space-time point of view can be regarded as a revolution in our conception of the physical world” (Geroch 1981, 33). This testimony alone should serve to distinguish the geometrized versions of “Newtonian” mechanics from canonical, non-geometrized, NM. Such a difference has not gone unnoticed by Malament (although he is here remarking on a slightly different problem): “I take it to be uncontroversial that [geometrized and non-geometrized NM] have the same domain of application, and have the same ‘observational’ or ‘experimental’ consequences (however one understands these notions). But the question arises whether they are equivalent according to

more stringent criteria” (Malament 1995, 508). Philosophical qualms being as they are, we might still ask if this difference is sufficient for the two to be disanalogous?

I see it as an uncontroversial empirical fact that most technical investigators (philosophers of science and scientists both) who view the reduction are nonplussed by the conceptual differences between α^* (geometrized NM) and β (non-geometrized NM). Recalling Hooker’s discussion of the analog relation from §2.1.2, I think that this serves as an excellent example of how “normative commitments” of scientists temper the judgment of an $[A_n]$. Here scientists are committed to allowing two characterizations to be “equal”, so long as they both describe the motion of objects adequately. This may belie some tendency to want to view science as an activity that “models the behavior of the world”. If it is the case that two models describe behavior identically, then they are each adequate, and may be considered equivalent. Conceptual differences between models are of little consequence, so long as each accomplishes the job of predicting motion. Here investigations about which model “really” describes the world are perhaps deemed extra-scientific, for empirically there cannot be a difference between the two.

Now let me attempt to summarize all that has been packed into our analog relation $[A_n]$, when employed in relating the geometrized NM (α^*) and non-geometrized NM (β). Any representation of dynamic motion captured by a β system may also be accurately described by an α^* system. Conversely, the procedure is a bit more nuanced. Take a subset of α^* spacetimes that have particular structural symmetries, called $\alpha^{*'}.$ Now, given an $\alpha^{*'}.$ spacetime, we may “recover” an infinite number of β -type systems that will also represent the motion of objects. Those are the mathematical features of $[A_n]$. Conceptually, $[A_n]$ recognizes that β and α^* are very different animals: the former has a potential field that moves objects about a Euclidean spacetime, whereas the latter has motion dictated solely by the structural features of a curved (non-Euclidean) spacetime. I think that in this particular example, each of these features exhibits the flexibility which an $[A_n]$ may possess. The described mathematical and conceptual differences seem rife, yet still we are willing to tolerate them and proclaim that NM has been reduced to GR. This is because for the typical GR practitioner, α^* and β provide equally effective representations of objects’ motions, making the analogy an acceptable one.

§2.2.1 Traditional Epistemic Goals – Theory

Explaining Theory:

Malament wrote a technical paper where he was very clear about one of the payouts the result provided: “insofar as it is the ‘classical limit’ of general relativity, Newtonian gravitational theory must posit that space is Euclidean” (1986b, 406). As was mentioned above, only certain classical spacetimes that are the result of limiting GR spacetimes may be recovered; specifically the ones that are spatially flat⁴³. This fact is interesting because it serves as an explanation: it tells us something about the class of classical spacetimes that result from the limiting process. Not all classical spacetimes are spatially flat; however it happens that the ones that limit from GR are. Malament saw this as one of the clear reasons that one would be interested in working through the reduction. Before the details of the reduction were worked through, one wouldn’t have seen spatial flatness as necessary – Malament mentions how this was a “big surprise” (1986b, 406). The epistemic payoff we see is rather straightforward as our *explanans* is our reducing theory and the *explanandum* is the reduced theory⁴⁴. Everything functions just as we would come to expect, according to the description given by Sarkar in §1.2.1.

Weatherall noticed another very interesting explanation that the reduction provides (2011). Newtonian theory has two notions of mass: an inertial mass, which is employed when we attempt to make an object move from rest, and a gravitational mass, which tells us how hard gravity “pulls” at an object. There is no help within NM concerning how these two compare to one another: it is simply an empirical question. Curiously, when we go to measure each of the two masses, it turns out that they are exactly equal. This is the *explanandum* which Weatherall focuses on: *why* are the two masses the same?

In GR, there really isn’t any notion of a “gravitational mass”, merely inertial mass – this is because there isn’t really any notion of a “gravitational force”. In the reduction that we have outlined above, when we apply the Trautman Recovery Theorem, the gravitational mass appears.

⁴³ Although *spacetime* is still be curved, by holding fixed the time and each spatial slice along the temporal dimension, we find the three dimensions of *space* are in every case rectilinear.

⁴⁴ This is a bit fast: the *explanans* is certainly α . However the *explanandum* is the intermediate theory, α^* , not the end reduced theory β . I have not made this distinction in the main body of the text, as I think it obscures the above point. If this is troublesome to the reader then one can read the reduction in question as being “from α to α^* ” for the purposes of arriving at the intended payout.

Furthermore, it turns out that “the coupling to the gravitational field in [NM] is given by the inertial mass”, thereby demonstrating that they are in fact the same quantity (Weatherall 2011, 432). Thus the gravitational mass “comes from” the inertial mass, and as such we have an explanation about why they are identical under NM. The *explanans* of the two-mass conundrum is thus GR. Again we see an excellent example of how a reduction can do significant epistemic work for questions otherwise unanswered by the reduced theory.

One worry might arise with regards to both of these explanations. Malament previews the allure of his “spatial-flatness thesis” with the following:

It is probably most natural to assume, and perhaps [Gauss, Riemann, Helmholtz, Poincaré] did assume, that any hypotheses about spatial geometry function only as inessential auxiliary hypotheses within Newtonian physics - superimposed, as it were, on a core of basic underlying physical principles which themselves are neutral with respect to spatial geometry. Yet it turns out that there is an interesting sense in which this is just not so, a sense which is only revealed when one considers Newtonian gravitational theory from the vantage point of general relativity. (Malament 1986b, 405)

Presume that one has burning questions about the structure of space and its necessity in Newtonian mechanics. Likewise we could have similar questions about the curious connection between gravitational mass and inertial mass. In each case we have been given an explanation, via the “vantage point” of Newtonian gravitational theory. Now imagine Newton, in the 17th century, being pestered by such questions about his theory. Would the following reduction be enough to satisfy him? We have perhaps good reason to think that he would recognize that – *in as much as it is limited from GR* – we have these answers. Yet this might not be enough, as Newton might not think that GR is at all of relevance here. I can think of one foundation for why our Newton could have this inclination, one that questions GR’s viability and worries about conflicting results from other theories.

One possible reason for this reaction could be that, in a 17th century context, GR could be seen as an unaccepted, controversial theory. We could even cast such a worry as “GR might not be true”, although I do not think it is absolutely necessary to do so⁴⁵. However construed, it seems quite relevant to think that, so long as GR is not considered a theoretic contender, how

⁴⁵ If it were an issue of truth, then it would invoke the discussion about whether false theories are able to serve a role as explainers. Philosophers of science, such as Alisa Bokulich, have already begun to question whether, in general, such a requirement is misguided (Suarez 2008, chap. 6). Weatherall also mentions this issue, likewise suggesting that forcing theories in the *explanans* to be true is entirely too stringent (2011, 428, fn 12).

NM relates to GR is rather uninformative. Spatially there are other, non-Euclidean, possibilities for theories of dynamics that are otherwise quite Newtonian. Just because space must be flat if we arrive at NM from the limit of some unacknowledged theory need not be any cause for interest.

The answer to the problem of thinking that the GR reduction does not in fact inform us of anything interesting about NM consists in reminding ourselves what kind of reduction we have provided. The case is an example of a successional reduction. This is just a roundabout way of saying that, in regards to the above concerns, it must be granted that GR is the current accepted theory, one that is both empirically adequate and scientifically robust. This assuages any 17th century worries, as it firmly acknowledges that we should care about how NM relates to GR, as GR is indeed no ordinary theory – it is NM’s successor.

Sarkar claims that epistemic explanations operate by one theory explaining the other, or features of the other. Focusing on successional reductions, Sklar claims that “an incorrect theory cannot be explained on any grounds” (1967, p.112). Sklar believes this because he is relying on a D-N account of explanation, in which there is no room for an “approximate” correlation between the succeeding and succeeded theory. Hopefully this section has provided two examples of how we can have a theory explain another in a successional reduction, *pace* Sklar.

§2.2.2 Differing Epistemic Goals – Reduction

Explaining Scientific Progress:

Sklar believes that there is room for explanations in successional reductions, just not of the sort described in §2.2.1. Instead, there is a “distinction between explaining a theory and explaining its apparent successes” (Sklar 1967, p.112). Recalling back to §2.1, this is also the sort of explanation that was evoked by the requirement of (III)(ii) in the Schaffner model, a requirement that was notably absent in the New Wave model of §2.1.1. This section will begin to demonstrate exactly how this other type of explanation can be achieved.

For the Schaffner Model, the requirement was that:

(III) β^* corrects β , in the sense that:

- (i) β^* provides more accurate predictions than β in most cases.
- (ii) β^* explains why β was incorrect, and why it was once considered correct.

Now this clause cannot be easily mapped on to our current case, as we found the New Wave model was much better fit to the reduction. The case study made modifications to GR, not NM, and as such we employed an α^* , not a β^* . If we were to ignore this difference and instead substitute our α^* (geometrized NM), we would arrive at the following:

(III_S*) *Geometrized NM* corrects *Non-Geometrized NM*, in the sense that:

- (i*) *Geometrized NM* provides more accurate predictions than *Non-Geometrized NM* in most cases.
- (ii*) *Geometrized NM* explains why *Non-Geometrized NM* was incorrect, and why it was once considered correct.

There is some issue whether this substitution is a legitimate one, as Schaffner conceived of the move as making an intra-theoretic comparison (between β and β^*), not the inter-theoretic move that we see above (between β and α^*). However, seeing as the instantiated criterion of (III*) still demands that the “analog theory” corrects the “reduced theory”, the substitution is at least plausible.

The work being done by (i*) is minimal. Once we have constrained α by limiting it to become α^* , the prediction of the motion of objects provided by α^* and β is the same. So it is false that, in these cases, geometrized NM fares better than non-geometrized NM. But we also know that the predictions of geometrized NM are still lacking: among other problems, there is not even an upper-bound for the speed of light. GR provides accurate predictions, yet by limiting λ we eliminate many of the features of GR that provide the improved accuracy.

Turning to (ii*), the requirement of having either GR or geometrized NM explain why non-geometrized NM was once considered correct might seem a bit strange. Similarly, the Suppes condition of explaining the “apparent successes” might suffer from the same associated triviality. We know from the onset that GR is the successor theory to NM. Before we even begin to worry about “reductions” or “fitting a model of reduction to the case”, this scientific fact holds because GR does a better job of predicting/explaining the phenomena than NM does. Inasmuch as the two theories are close to one another, NM was successful for simple reasons: its

predictions are close to the predictions of the better theory, GR. Why was NM successful for so long? Simply because GR hadn't been devised yet, making NM the preferred theory for some time. Now, if we are to require that our reduction *further* tell us why NM was successful, the requirement appears overly demanding or, as put by Larry Laudan, “clearly *gratuitous*” (1981, p.43, emphasis in original).

I think that there is additional epistemic work being done in this case however. The limiting of $\lambda \rightarrow 0$ on $(\mathbf{M}, \mathbf{g}_{ab}(\lambda))$ had the effect of opening the light cones at any given point. The limit has the “maximal particle speeds go to infinity” (Malament 1986b, 406). This is something which is disallowed in GR, but in (geometrized and non-geometrized) NM there is no bound on the speed of particles. Thus we can begin to see the following story take place: if we don't allow there to be a “maximum speed”, then the picture provided by GR will resemble the one provided by NM. And this has an effect of explaining NM's successes, so long as we recognize that scientists at the time had no compelling reasons, experimental or theoretical, to presume that there was a bound on the speed of objects.

During Newton's time, terrestrial experiments were simply incapable of reaching speeds close to the maximum speed that GR now recognizes. The difficulties of even breaking the sound barrier were rife. Even the movements of bodies such as planets, or electromagnetic phenomena within wires, made any prediction of an upper-bound for speeds difficult. Much of this was a reflection of the technical problems associated with measuring such high speeds, including the lack of precise measuring apparatuses. There are some cases where the predictions provided by GR and NM differ wildly; in these cases the need for explanation is certainly not “gratuitous”. But these cases simply were not common occurrences for the majority of the time during which NM was considered the leading physical theory. These reasons provide a richer description of why NM was so successful.

Now that §2.2 has shown us details, it is clear that a story along the lines of “Newton didn't have the same mathematical tools that Einstein did” won't work as an explanation. Although this is true, the formulation of a geometrized version of NM could have instead been used. Here using this theory, really just our α^* , wouldn't affect the predictions of NM in any way, as α^* and β are equally effective in this respect. Instead it is not the mathematics, as much as the assumptions about motion, which caused Newton to “get it wrong”. These assumptions are the significant cause of the difference between the two theories, and now we see how Newton

and others were not unjustified in making them, given the historical context from which they were operating. Thus I think that the right idea was in place by the requirements of (III), although focusing entirely on α^* is perhaps misrepresentative. Instead, it is the entire transformation of GR into geometrized NM – in addition to the historical context surrounding each theory – which facilitates the discussed explanations.

I think that these are some of the more interesting lessons to be learned from the GR/NM reduction. There is something very valuable about learning why a theory was successful, and the process of reduction can inform us of important factors. Sklar, Wimsatt, and Schaffner all realized that this had some place in the discussion and modeling of reductions. Next in §3, we will turn to the specifics of a similar reduction to demonstrate the potent epistemic goal of having the succeeding theory “explain away” the succeeded theory.

§2.3 Conclusion:

§2.1 showed attempts to create models of reduction that allowed the reduced theory to err. The first model was Schaffner’s, involving an intermediary theory created from the conceptual resources of the reduced theory. This intermediary theory is analogously related to the reducing theory, in that it will mimic its structure and conceptual underpinnings. Schaffner stipulated other requirements for a reduction, notably the provision that the “[analog theory] indicated why [the original reduced theory] worked as well as it did historically” (1974, 618). In this way he was able to create a picture of reduction that allowed for *approximation*, as it allows that the original, reduced theory may be “close” but imperfect in its description of the world.

§2.1.1 looked at a revision to Schaffner’s model. The New Wave challenged the assumption that the analog theory need be created from the resources of the reduced theory. Instead the New Wave required that the analog theory be generated using the reducing theory’s concepts and vocabulary. Additional requirements, such as the stipulation that the analog theory explain the historical successes of the reduced theory were dropped.

The chapter’s case study was the reduction of NM to GR. The details of this case were fit to the New Wave model, explicitly showing a scientific case where a reduction employs an intermediate theory constructed through the conceptual resources of the reducing theory. After

the reduction was explained in §2.2, §2.2.1 showed how the reduction accomplishes traditional epistemic goals of explaining features of NM by GR. First, the reduction shows how space becomes Euclidian once we limit GR spacetimes. Secondly, the correlation of the inertial mass and gravitational mass ceases to become an empirical curiosity, as the reduction demonstrates that the two masses are “split” from GR’s singular conception of mass.

Finally in §2.2.2, it is shown that Schaffner’s goal of “explaining prior historical successes” can still be a function of the reduction. The reduction demonstrates what differed *conceptually* in NM from GR, and then showed exactly what these differences entail when examining the two theories. Notably, it facilitates an understanding of why NM was unchallenged for so long, and how GR challenged NM only after centuries of prominence. The reduction illustrates how presumptions made by NM about motion, space, and time explain its past successes.

Chapter 3 - Limits and Approximations: Special

Relativity to Classical Mechanics

We all know what “approximately” means, but just try to say what it means for the general case.

Lawrence Sklar (1967, 111)

§3.0 Introduction:

The first philosopher to seriously attempt to incorporate limiting processes in reductions was Thomas Nickles (1973): §3.1 overviews his distinction between a “philosopher’s reduction” and a “physicist’s reduction”. After discussing Nickles’s account, we will turn to limiting the momentum equation of special relativity (SR) to yield the momentum equation employed in classical mechanics (CM)⁴⁶. This case study is often cited as a quintessential reduction by philosophers and physicists alike (Batterman 2002a) (Hooker 1981a) (Rivadulla 2004). There are three ways that the literature has examined the limit: by focusing on (i) the speed of light (c), (ii) the velocity (v), or (iii) the quantity $(v/c)^2$. §3.2 focuses on what limiting c would entail for the momentum equation of SR. In his discussion of the result, Nickles cautions against the limiting of c in the equation, and concludes that we should instead focus on limiting v . I claim that it is fine to allow c to vary, arguing that the criterion we must employ to justify limiting processes in physics should be grounded in the physical and conceptual meaning that is attached to the limit of a value, rather than merely a terminological distinction. By doing so, I claim that much of the insight gained by the limiting relation follows from this important justificatory process. This prompts a discussion in §3.2.1 of what goals underlie such a process: they are indeed explanatory, but not in the traditional sense of one theory explaining another. Instead they explain features of the successes of CM, in a way similar to those we observed in §2.2.2.

⁴⁶ In some ways, the SR-CM case might be seen as a special case of the GR-NM reduction, however the details of how the reduction plays out, as well as the philosophical literature which has discussed these details, are in both cases different. Thus I have chosen the labels of “SR”-“CM” and “GR”-“NM” to distinguish which reduction I am referencing, not to essentially indicate any larger conceptual difference that might exist between the two cases.

In §3.3 we turn to v , again looking at the mathematics of the result in detail. I show that the limiting relation for v , abstractly, is quite weak, as a large number of curves also possess such a property. In light of this observation, I argue for why the result carries so much significance – for philosophers and physicists – in the first place: only when we include contextually relevant features, involving details both historical and experimental, can we distinguish the importance that the SR and CM equations have from other, less interesting, functions. There are two significant goals that come from our discussion of v : §3.3.1 again shows how extra-theoretic explanations may be achieved, while §3.3.2 details how an older theory may transfer confidence to the newer theory.

Lastly, Fritz Rohrlich and Larry Hardin describe the result by taking the limit of $(v/c)^2$ (Rohrlich & Hardin 1983); §3.4 looks at this approach, as well as a similar one presented by Robert Batterman. Rohrlich and Hardin articulate what it would mean for a theory to be “established”, a device intended to showcase the predecessor’s role in the development of science. I analyze this concept in §3.4.1, finding it tangential to our purposes of theory comparison. §3.4.2 shows how “domain of validity” of CM relative to SR allows us to succinctly restate an explanation of the successes of CM that were highlighted in §3.2.1 and §3.3.1.

§3.1 Nickles’s Two Models for Reduction:

Nickles distinguished between what he saw as two separate usages of the term “reduction”, what he labeled *reduction₁* and *reduction₂*. His “reduction₁” was the “philosopher’s reduction”, one that provided an ontological or conceptual economy of entities or terminology. Writing in 1973, Nickles saw this as the predominant archetype focused on by philosophers in the literature; he took both Nagel’s derivational model (as discussed in §1.1.1) and “the reduction of optics to electromagnetic theory” to qualify as examples of a reduction₁. Although Nagelian reduction has since been modified in many ways, and the optical result is in need of further scrutiny⁴⁷, the overall gist of a reduction₁ should be clear. A reduction₁ from one theory to another occurs when there is a conceptual or ontological consolidation, or a general increase in

⁴⁷ Let me suggest that, without supplying any details, this is perhaps a good example of Schaffner’s “Cheshire Cat” problem – a reduction that is at best creeping, not sweeping (2006). Even at the time of his paper, Nickles is well aware of the controversies that surround many of the “reductions” he provides as examples (1973, footnote 1).

efficiency in the organization or presentation of a theory (Batterman 2002a, 183). With regards to our brief classifications of goals for reduction in §1.2.1, the motivation behind a reduction₁ is typically ontological or epistemic. The other important feature about a reduction₁ is that it goes in the “natural” direction: α is a “more fundamental” theory (such as neurobiology) that is used to arrive at a reduced theory β (such as psychology). Thus we read the relation as “reducing β to α ”.

Nickles’s major point in his paper was not to delimit what constituted a reduction₁ specifically, but instead to distinguish it from another, then less-recognized sense of “reduction”. His “reduction₂” was a “physicist’s reduction”. Nickles notes that “the great importance of reduction₂ lies in its heuristic and justificatory roles in science” (1973, 181). Here the purpose behind a reduction₂ need not focus on the explanation of one theory by another, but instead tells a tacit story of scientific progress, much in the way we saw in §2.2.2. Such a reduction is typically employed as a successional reduction between a predecessor theory and its successor; a reduction₂ helps scientists to justify a new theory if it was seen to reduce₂ to an older theory in the appropriate regime. In a reduction₂ we begin with the successor theory and then perform “limiting processes and approximations of many kinds” so as to arrive at the predecessor (Nickles 1973, 183). In doing so notice that the order of the theories is switched: by limiting the successor theory α to the previous theory β , we find that we will “reduce α to β ”. In this manner we find that the new theory “looks like” the old, previously accepted theory, within the limit of certain parameters. This is informative to scientists, as it transfers the confidence they once had for the old theory to the new theory, at least in the accepted boundaries, as §3.3.2 will elaborate.

We examined the GR-NM case in §2.2. This clearly is a successional “physicist’s reduction” – for Nickles a reduction₂ that is to be read as “reducing GR to NM”. Yet the New Wave model, which was quite a good fit to this case, claims that we had in fact been “reducing NM to GR”. So who is correct? I think that each has a place. The New Wave, seeing that a reduction tells us which theory is more primary, is correct in reading the relation as they have, for GR *is* the successor of NM. However a physicist, paying attention to where we start (GR), and where we end up (NM), will describe the process as “reducing GR to NM”, as it is a good descriptor of the reduction process. So I do not see a significant conflict between the two locutions: each is an adequate descriptor of how to read the direction of “reduction” in different academic contexts.

Does a reduction₂ apply to cases that past models of reduction we have examined cannot capture? Sarkar believes so, as he claims that for the cases that a reduction₂ seeks to describe, “there is clearly no question of the derivation of α from β . Thus, a Nagel-Schaffner type of theory reduction cannot be obtained” (Sarkar 1992, p.173-174). Certainly Nagel does not allow for any approximate difference between α and β , so under his model this is correct. Indeed Nickles sees Nagel’s model of reduction as exemplifying a reduction₁ – the very motivation for him to devise an alternate account by his reduction₂. Under Schaffner’s model this is debatable. Nickles himself believes that Schaffner’s model cannot incorporate limiting, claiming that limiting would have to occur in the derivation of β^* . This seems quite difficult, as we must arrive at β^* only by means of α and bridge laws. Another way of trying to construe the Schaffner model as a reduction₂ would be to subsume the limiting operation within $[A_s]$. The trouble with this suggestion is that $[A_s]$ is intra-theoretic, and yet it seems more natural to view the operation of “limiting one theory to another” as occurring inter-theoretically. If any of the outlined models may be said to be a reduction₂, I think that the New Wave model is the best fit. As discussed in §2.1.1, limiting can be considered part of **C**, the conditions imposed upon α . Indeed it was the intention of Hooker, Churchland, and others for their model to be able to account for limiting. Thus if we take the New Wave to be representative of a “Nagel-Schaffner type” reduction, then Sarkar’s statement is incorrect. §2.2 provides a reduction that is well-fit to the New Wave model, yet also seems to be a paradigmatic successional reduction that involves limiting, i.e. a reduction₂.

Although Nickles does not specify the particulars of a reduction₂, Batterman attempts his own rendition (2002, 18, 78). Batterman envisions the “physicist’s reduction” to exclusively involve limiting:

$$\text{Lim}_{\varepsilon \rightarrow 0} \alpha(\varepsilon) = \beta \quad [3.1]$$

Here we limit one or more parameters that are present in α so that we may arrive at β . Batterman intended this equation to be employed primarily to theory parts, such as specific equations of a theory, and we shall employ it below to just the momentum equation. This is fine, as it rarely would make sense to “limit an entire theory” or something to that effect. It is still questionable, however, whether [3.1] is an accurate representation of Nickles’s criterion, because Nickles is

willing to allow that we may employ limiting processes yet also “approximations of many kinds” to α (Batterman 2002a, 183). It is then appropriate to consider Batterman’s [3.1] to be an explication of one type of the processes that Nickles allows. Batterman’s take on the “physicist’s reduction” is appropriate to begin our discussion, because [3.1] leads straightforwardly towards our main question: what sort of objects can be represented by ϵ ? When is limiting ϵ permitted?

§3.2 The Problem of Limiting c :

Nickles makes reference to the limiting of the SR equation for momentum to the CM equation, a case that he sees as “epitomizing” the overall reduction between SR and CM⁴⁸ and in my mind a quintessential case of a reduction₂ relationship. Here $\mathbf{p}$ represents an object’s momentum, $\mathbf{m}$ is the mass of an object⁴⁹, $\mathbf{v}$ is its velocity, and $\mathbf{c}$ is the speed of light in a vacuum:

$$\mathbf{p} = \mathbf{mv}/\sqrt{1-(\mathbf{v}/\mathbf{c})^2} \quad [3.2]$$

When discussed by other philosophers of science (Batterman 2002a), (Hooker 1981a), and (Rohrlich & Hardin 1983), typically $\mathbf{c}$, $\mathbf{v}$, or $(\mathbf{v}/\mathbf{c})^2$, are limited to ∞ , $\mathbf{0}$, or $\mathbf{0}$, respectively, to arrive at a representation of the familiar dynamic equation for momentum from CM:

$$\mathbf{p} = \mathbf{mv} \quad [3.3]$$

Regardless of differences in approach, the core of the argument is undoubtedly the same in each case: the mathematical limit of one function will yield the other function, just as described by [3.1]. Often a passing phrase is provided for the envisioned physical implications of the limit. For example, $(\mathbf{v}/\mathbf{c})^2 \rightarrow \mathbf{0}$ may carry with it a description of “the limiting domain where velocities are small compared to the speed of light” (Batterman 1995, 172) or a reference to a realm where “objects move with velocities that are small compared with the value of $\mathbf{c}$ in empty space”

⁴⁸ In what follows we will only be dealing with operations to the momentum equations of both SR and CM. For the sake of brevity, I will refer to “reducing SR to CM” rather than the longer, more laborious “reducing the momentum equation of SR to the momentum equation of CM”, despite the latter being more accurate.

⁴⁹ Nickles employs “rest mass”, $\mathbf{m}_0$, in both the relativistic and classical equations. As was previously recognized in §1.1.1, this is unnecessary (Rivadulla 2004).

(Rivadulla 2004, 418). Nickles himself refers to the $v \rightarrow 0$ limit as “the limit of low velocities” (1973, p.184). Sometimes authors will justify why they have chosen to limit one parameter rather than another⁵⁰, but such discussion is rare.

For the purposes of this section, we will focus our discussion on limiting the speed of light, c . When we limit c to ∞ , the term v/c will limit to 0 , leaving [3.3] very cleanly after the expression is cleaned up. Regardless of how smooth such a result may appear, Nickles is a bit cautious of such a maneuver and wary of allowing c to vary, a process that must occur when taking a limit. He warns that we must remember that c is a *physical constant* in the equation: a stand-in for a very specific number that has been experimentally obtained. Performing a limit on a constant is mathematically illicit, essentially by definition. The number represented by c is a fixture of the equation as much as the operations signified by the division or squareroot signs. To help show his point, Nickles presents two polynomial equations:

$$w = ax + 2y + g$$

$$z = bx + ey + d$$

Here x, y, w, z are “physical variables”, and a, b, e, g, d are “numerical constants” (Nickles 1973, p.198). Nickles’s worry is that by limiting the constants, one can trivialize the differences between any two theories. For instance, if $a = b$, and $e = 2$ we could limit g to d and then the two equations would be the same. We could limit a and b to 0 , let e approach 2 , and then by limiting d to g we would again have the two equations reduce. The situation gets much worse:

...by letting $2 \rightarrow 0$ we can also eliminate the y term. And by this means every equation reduces to every other – a complete trivialization of the concept of intertheoretic reduction. Any physical-constant coefficient can be eliminated by taking it to 1 (take additive factors to 0). Any expression whatever may be introduced or eliminated from an equation by these means. Clearly we must say that letting numerical constants change value is mathematically illegitimate. (Nickles 1973, p.199)

In an attempt to salvage reductions from triviality, Nickles concludes that limiting “constants of nature” such as c is illicit in most cases.

⁵⁰ For example Batterman argues that choosing $(v/c)^2$ is most appropriate, stating that “it is best to think of this limit as $(v/c)^2 \rightarrow 0$ rather than as $v \rightarrow 0$ since $(v/c)^2$ is a dimensionless quantity, so this limit will not depend on the units used to measure the velocity” (2002a, 79, fn. 3). Justification for each way of conducting the limit will appear in appropriate sections: §3.2 for limiting c , §3.3 for limiting v , and §3.4 for limiting $(v/c)^2$.

Nickles is however well aware that theoreticians often *do* limit constants, despite his apprehension (he cites Planck and Bohr as being on his list of offenders). Nickles also recognizes the benefit that may at times be garnered by “tinkering” with theoretical stalwarts such as Planck’s constant (Batterman 2002a, 201). Indeed, we saw Trautman and Malament “limit c to ∞ ” in §2.2. Nickles cautions that we should require the transformations, such as limits, that surround reductions “make physical sense as well as mathematical sense” (Batterman 2002a, 201). The latter has been explained, and as an example of the former he posits that any case which involved limiting the temperature to infinity in any equation would be nonsensical (Nickles 1973, p.200). Is limiting c to infinity likewise nonsensical? Or can we provide some justification, physical, mathematical or otherwise, for this “tinkering”?

To begin the argument for how we are to go about reducing SR to CM by varying c , I will provide a naïve reconstruction of how to construe “limiting of c to ∞ ”. What follows is the story that first occurred to me when I considered what such a limit might entail: “Take all possible worlds that behave according to [3.2]. Order them according to differing values of c ; c is a physical constant, and may well be measured differently in each case. Now we can certainly observe the effect of the increase of c in such an arrangement: as c becomes larger, we can question the effect it has on various physical relationships. It turns out that, the bigger c gets, the more [3.2] approaches [3.3]. Indeed, take a value of v for which the desired difference in p between [3.2] and [3.3] is ϵ . There will be a possible world somewhere on the array of possible c -worlds that will put this p value within ϵ for that v ”⁵¹. By proceeding in this fashion I think we may get a very reasonable picture of how one might come to understand the process of “limiting the constant c ”.

I wish to claim that such a story of ordering possible worlds according to various c -speeds is how we may begin to understand how taking such a limit is reasonable. Such a story is counterfactual, but this should not provide any great worry in itself (Sarkar 1992, p.174). Indeed when limiting c in a situation much like the reduction described by §2.2, Michael Redhead refers to the manipulation of c as a “counterfactual variation” (2004, 528). At a glance it may appear that in the above we are thus treating c as a variable (as we permit it to vary). However in each possible world c is a constant. So should we deem it a variable, and thus seemingly disobey

⁵¹ Other facts concerning convergence must naturally obtain, but I have here been purposefully imprecise as I do not think that the absence of such details will obscure my point.

mathematical rules that Nickles is concerned with, or as a constant? In general, Nickles's apprehension seems well-founded, for limiting any parameter we see fit can surely lead to problems. But there are still instances where such limiting is warranted, mathematically and philosophically. This is especially true in our case because, as we shall see, there is such a lucrative philosophical and scientific pay-off.

The lesson to be drawn from our story is that we should disavow any reliance on a rigid determination of what demarks a “variable” and what demarks a “constant”. Linguistically such a distinction may do useful work, but I doubt that any abstract criterion can be proffered so as to give guidance outside of an equation's context. Take [3.2] as it is used in a simple high school physics class – here c functions as a constant⁵². But in the described context of our possible-world discussion, it seems much more appropriate to refer to it as a variable. I am not looking to lecture mathematicians/physicists about how to employ their terminology; instead my intent is to show that however these lines are drawn, a “variable/constant” dichotomy gives no guidance when determining when limiting operations are permissible.

So what must obtain for us to be permitted to limit a value⁵³? My claim is similar to the criterion offered by Nickles himself: only when a limit has some conceptual/physical meaning, not merely when we know that it will provide the desired result. For instance, I could decide to limit v/c to 1 in [3.2], and here this would be justified as “letting the speed of an object approach the speed of light”. I would seek to claim that in this case “variable” limiting is allowed because it represents an understandable physical process. By the same criterion limiting m to c would not be justified, as there is no well-demarcated physical situation that it could possibly represent (just look at the units). Naturally examples such as limiting the number 2 to the number 3 could never be justified, as such a story would never make mathematical sense. However Nickles's example from the passage above does not function as a *reductio ad absurdum* for any general sort of argument. Such a case is simply disanalogous to ones like that of c , because limiting a numeral in some analytic discussion is quite different than talking about the behavior of physically-determined constants. In the passage quoted above, Nickles uses “numerical constants” and “physical constants” interchangeably – this is a mistake, for the two operate very differently. To

⁵² In other contexts, the speed of light is regulated by the medium through which it travels, and thus in this respect may be said to vary. Indeed, to a more sophisticated relativity practitioner, most talk of a “speed of light” is misguided.

⁵³ I here am using “value” in a hope to avoid any limit-related biases that might instead be lingering were I instead to use “constant” or “variable”.

return to our discussion of c , I believe that my fanciful “continuum of c -worlds” description above is part of the counterfactual justificatory story that allows us to limit c to ∞ in [3.2]⁵⁴. We must say more, however. The above story merely shows that to speak of c increasing without bound is conceptually reasonable. To complete the story however, we must answer the following looming question: what would it mean for the speed of light to be “infinite”, given the context of [3.2] and [3.3]? By providing an answer, we will be able to understand what the result is ultimately meant to tell us.

We have a fairly good idea of what the measured speed of light is: by recent measurements, approximately 299,792 km/s (Bortfeldt 1992, 3–38). Even back in Newton’s time, observations made by Jean-Dominique Cassini and Ole Rømer revealed that the speed of light appeared to be finite. Calculations based on these data made by their contemporary Christian Huygens placed the value somewhere in the vicinity of 220,000 km/s (Bobis and Lequeux 2008). So what would “letting c limit to ∞ ” represent, if it is quite clear to all parties involved that c is bounded? Huygens did his calculation in 1690. At the time when he revealed his result, Robert Hooke remarked that c was “so exceeding swift that ’tis beyond Imagination... if so, why it may not be as well instantaneous I know no reason” (Daukantas 2009, 46). I think that this reaction is quite telling. Hooke, yielding to pragmatic impulses, realized that such speeds were virtually unobtainable for any earthbound experiment. Indeed the observations of Rømer were conducted during the eclipse of Io by Jupiter. Huge distances were needed to make the basic determination of whether or not light had a finite speed at all. (Kepler earlier had predicted that light moved over distances instantaneously; however this was grounded more on *a priori* reasoning). Isaac Newton himself in *Opticks* recognizes the observations of Rømer, tacitly conceding the finiteness of light’s speed (Newton 1718, 252). But such an acknowledgement, on its own, tells us little.

The important point, in my mind, is still Hooke’s. One can begin to see the argument for why Newton, when making his theory of dynamics, was not in any way influenced by a finite speed of light. Instead his theories were concocted by assuming that light moved so fast that it perhaps “may not be as well instantaneous”, to requote Hooke. This is why we would like to take

⁵⁴ Nickles very briefly refers to cases where we could “think of [Planck’s constant] and c as variables in a noncommittal metalanguage” (1973, 201). This cryptic remark is referenced by a footnote telling how we may limit constants so long as we regard them as “variables in the metatheory” (Nickles 1973, 201). Charitably, I think that this may well be referencing a procedure much like the story I have provided above. Lacking further description from Nickles, however, I am hesitant to fully ascribe this position to his paper.

a limit of c to ∞ , because in doing so we will show this: that if Einstein were operating under assumptions similar to those of Newton, then the two would have come to a closer understanding about how momentum worked. These thinkers arrived at different equations, but the limiting of c shows that, in a certain sense, they differ only inasmuch as they relied upon different presuppositions⁵⁵. It informs us how, in Newton's day, it essentially would not matter that the speed of light was finite, for such a value was so large that the empirical implications would seemingly remain untouched. For any calculation of momentum that a practicing scientist such as Hooke would make, the values given by [3.2] and [3.3] would differ very little. By giving a justification of why such a limit is reasonable, we are lead to the conceptual payoff, to the story that begins to tell us why such a limit is interesting to discuss in the first place. Allowing c to increase without bound tells informed scientists a surprisingly rich fact: that Newton and Einstein both came very close to capturing the behavior of objects' momentum, *modulo* the nature of light.

For us to be able to take the limit of a value, I claim that it must be justified in that it makes physical (and conceptual) sense. In this manner we may recognize how one could even be allowed to limit a physical constant, such as c . I think that there is a tacit story that lurks behind the scenes of the c -limit, and that the conceptual payoff of the result is larger than a mere mathematical relation. There were two significant components to the justificatory story I employed for why we can limit c to ∞ . First there was the "possible c -worlds" array that gave us confidence in the mathematical maneuver of operating on c . Second, there was a contextual component about what the speed of light was assumed to be at the time the theories were conceived, and how this related to experimental capabilities in the past several centuries. Aside from doing work to justify the mathematics employed, these components are also the key to understanding the result; without them we are left with a discussion confined to pure mathematics that is, in the abstract, quite shallow and uninteresting.

⁵⁵ To be clear, that the equations are similar tells us nothing about the similarity of the theories that they are embedded in; conceptually, Einstein and Newton entertained very different understandings of dynamics.

§3.2.1 Explaining Scientific Progress

When there is an epistemic goal operating behind a reduction, it is often assumed that it will be some sort of explanation of the reduced theory by the reducing theory. This was not what we saw occur when limiting **c**. This is a fact that did not go unnoticed by Nickles. He saw one of the relevant distinguishing features of a reduction₂ was that it did “not involve the theoretical explanation of one theory by another. Not all reduction is explanation!” (Nickles 1973, 185). Although this statement by Nickles seems correct, it is in need of clarification. Granted that it may well be that none of the work in the result can be described as “providing an explanation of CM by SR”, it clearly appears that the reduction’s goal was to provide an explanation, just of a different sort. We understand why CM was successful and unchallenged for so long. We understand how SR was slow to be accepted and why it was also a legitimate scientific challenge to CM. We understand how the one might be seen as “progress” in relation to the other. Each of these is an explanation. However none of them are an instance of one theory explaining another theory or portion thereof. And so we should conclude that there are epistemic payoffs to reductions that are not trans-theoretic explanations. Instead, in each case we find that the *explanandum* is something extra-theoretical; the *explanans* is the reduction result itself, along with the above stressed contextual details. The traditional models of explanation in reduction that would place **α** as the *explanans* and **β** as the *explanandum* (Sarkar in §1.2.1) just don’t fit this case.

This is not to say that there can never be theory-to-theory explanations in a reduction₂. I claimed above that the GR/NM reduction should be considered a reduction₂. §2.2.1 showed two examples of how GR explained otherwise primitive components of NM. Thus even in the intended sense of explanation, Nickles is incorrect to claim that “reduction₂ does not involve the theoretical explanation of one theory by another” (1973, 185). The only outlet is to claim the discussed case of §2.2 isn’t a reduction₂, but rather a reduction₁ – a position that is quite difficult to argue.

The explanatory payoff referenced by Nickles is quite similar to the type discussed in §1.2.1, the one mentioned by Sklar and Wimsatt in their discussions of successional reductions, and also the condition of (III)(ii) in Schaffner’s model from §2.1. The limit allows us to explain

why the CM momentum equation was successful for so long, and it gives us a hint of where CM is deficient. In addition, we have a straightforward reason why this reduction continues to get attention in introductory physics books and science lectures⁵⁶. The reduction, once completed, contextually embedded and understood, will explain facts that are relevant to both SR and CM. It provides a rich relational and scientific story that, without the details supplied by the various “limiting” relations, would otherwise be difficult to achieve. To make a fast analogy: popular science often describes modern science in a manner that obfuscates detailed technical matters. Here is one: “Take an ordinary tennis ball and throw it against a smooth concrete wall so that it bounces back to you...What if, on one occasion, the ball passed right through the cement wall? And what if it was only a matter of percentages? Fifty-five times the ball bounces back to you; forty-five times it passes through the wall!” (Lederman and Hill, n.d., 26). This description of quantum tunneling is unsatisfactory because the technical details of particles being able to overcome potential barriers are obscured by suspect macroscopic analogs, yet these details cannot otherwise be communicated easily to those interested. Similarly, a lot of the claims that surround SR/CM may likewise be opaque: “Einstein improved on Newton’s theory” or “after hundreds of years the accepted dynamics equations were proven correct”. Now the limit, along with some relevant contextual and historical details, provides the desired explanation in a suitably technical fashion.

§3.3 The Problem of Limiting v :

Now we turn to limiting v . Of the various parameters that could be limited, choosing to limit v was the original suggestion of Nickles, since he was suspicious of limiting c . First, notice that the process employed in limiting [3.2] to [3.3] with respect to v is not a simple single-variable function being limited to a value. Instead, we are limiting a multivariable function to another function. The important difference for our purposes is that we are not speaking of convergence to a point but instead a convergence of one curve to another. When Nickles mentions the result, he cites that [3.2] reduces to [3.3] when we let $v \rightarrow 0$ (Batterman 2002a, 182). Look at [3.2]: as v

⁵⁶ Of course it may well be that there is the didactic reason of providing a scientific example with which to garner more practice for mathematical techniques. However such instances find the relevant discussion relegated to the “exercise” section, and there are unexplained.

approaches 0 , the equation also approaches 0 . In fact, outside of context were we to offer such a problem to a beginning calculus student, they would quickly decide that the limit as v approaches 0 in [3.2] is 0 ⁵⁷ and move on to a more difficult problem. Instead, the move I envision is one like this: as v gets closer to 0 , [3.2] will behave more and more like [3.3]. The distance between the two functions, $|[3.2] - [3.3]|$, also vanishes as $v \rightarrow 0$ ⁵⁸. Were we to ask the student to be more rigorous, I submit that she would state something akin to the following: “tell me how close you would like the two functions to be to one another. Call this measure ϵ_1 . Then, *ceteris paribus*, there exists a v_1 , with $0 < v_1 < c$, such that [3.2] and [3.3] will be within ϵ_1 . Furthermore, if you were to desire the two to be even closer to one another, say $\epsilon_2 < \epsilon_1$, there exists a v_2 , with $0 < v_2 < v_1$, where again we will find that [3.2] and [3.3] are within ϵ_2 ”⁵⁹.

Notice how weak the aforementioned “limiting” relation now appears. So long as any two continuous curves agree upon the value that they are being limited to, they may be said to limit to one another as they approach that value almost regardless of their behavior at other points. The SR equation provided in [3.2] could have instead taken the form:

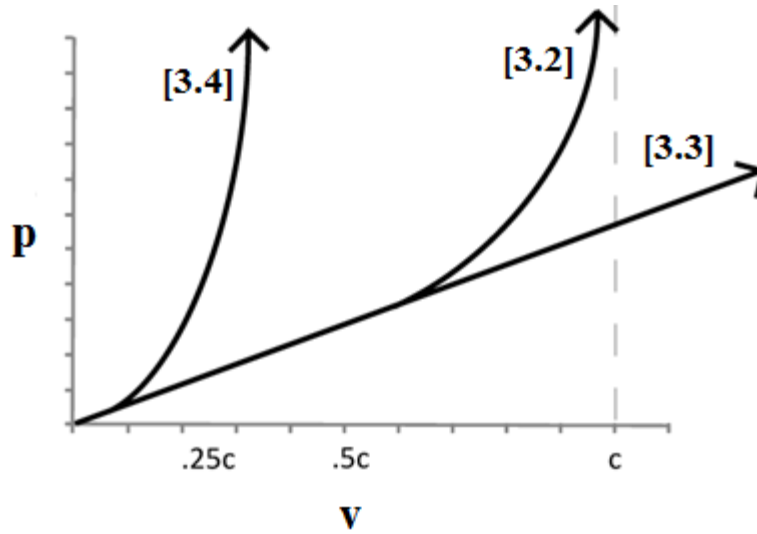
$$p = mv(1+v/c)^{50,000} \quad [3.4]$$

As $v \rightarrow 0$, we will still have [3.4] limit to [3.3]. By the criterion of [3.1], this would thereby be an instance of a physicist’s reduction. One could even qualify the result so as to mention that “in the limit of low velocities” [3.4] reduces to [3.3] (Batterman 2002a, 182). Refer to {Figure 3.1} to see each of these curves graphically limit to the origin.

⁵⁷ Analytically, you may substitute $v=0$ into [3.2]. Here, so long as we observe that c and m are positive, the equation behaves as “ $0/1$ ” and thus limits to 0 .

⁵⁸ This is not an assurance of a limit, as is elaborated upon below. I mention this as another indicator of the behavior of each function, not as evidence sufficient to show that the two have any further mathematical relationship.

⁵⁹ This is not an oblique asymptote, for we are not limiting to infinity. The difference between the equations does limit to 0 as they approach 0 . But when limiting to a value, as the next paragraph will show, the fact that the difference between two functions disappears says very little.



{Figure 3.1}

One possible objection at this point is that [3.4] is in no way a reasonable physical equation for momentum. Nickles brings up such a concern, and as a remedy he suggests that we leave our discussion of reduction_2 to concern only theories that are “established” (Batterman 2002a, 182). Here theories that are not legitimate scientific contenders, such as [3.4], can be dismissed out of hand regardless of whatever mathematical limiting behavior they exhibit. We might be able to corral such an objection, for despite [3.4] being a ludicrous physical theory, the fact that “[3.4] does still limit to [3.3]” helps us understand that any “established” theory will limit to [3.3] so long as it is: (1) continuous and (2) has $\mathbf{p}=\mathbf{0}$ when $\mathbf{v}=\mathbf{0}$. Here, even granting that only serious scientific theories are permissible in discussions of reduction_2 , we still are left with a too-easily satisfied property. It appears to be a straightforward necessary condition that for any momentum equation to be considered a serious scientific theory, it would require $\mathbf{p}=\mathbf{0}$ when $\mathbf{v}=\mathbf{0}$ ⁶⁰. Continuity is less-easily granted, but it seems difficult to avoid in the macroscopic regime. This would now force every possible “established” momentum equation to possess the property of “reducing₂ to [3.3]” – thereby damning the reduction_2 relation to triviality for the case of limiting $\mathbf{v}$.

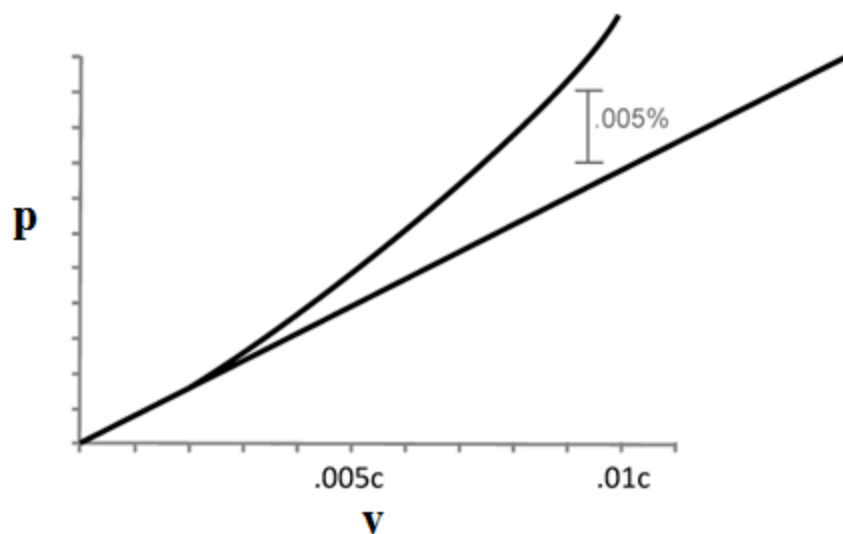
Given the large number of equations/curves that could be said to limit to [3.3], what would make the result described by Nickles, and many subsequent authors, interesting? *Any* random theory, established or not, that has a continuous formula for momentum that was

⁶⁰ I will actually dispute this claim in §3.3.1, but here leave it as intuitive and unchallenged.

dependent on v and went to 0 as $v \rightarrow 0$ would reduce₂ to [3]. Was Nickles just wrong about the reduction₂ of v , or can we still resuscitate the significant features of his “physicist’s reduction” for this case? I think there is a story to be told that salvages the importance of the result, and distinguishes [3.2] from [3.4], but that does not rely on the ineffective “established theory” criterion. My suggestion is that we rely on *context*: more specifically, on the empirical consequences of the two theories and how they interface with a historical progression of science. Bringing in such factors will allow us to see the reason why the result receives so much mention by scientists and philosophers of science alike.

§3.3.1 Success and Context

Reexamine the graph of [3.2] and [3.3] in {Figure 3.1}. As v approaches c , the two curves differ drastically: the SR equation asymptotes to infinity, whereas the CM equation remains linear. And, what is the important cited result, as $v \rightarrow 0$ the relativistic equation approaches the classical curve. To distinguish the SR equation from the many other possible and wildly aberrant equations (recall [3.4]) that also possess this property, I want to examine not *that* it approaches but *where*.



{Figure 3.2}

Look at how close together the functions are when $0 < v < 0.01c$, as we can see on {Figure 3.2}. At worst, the two equations differ by .005%. Without other details, this information says very little: it says that if we look at 1% of the graph, this narrow area puts the two curves “close” together. I use scare quotes because again, devoid of context, .005% can either be very close or very far off. For example, if I were to measure that the gravitational constant was .005% off of the established value, and my equipment/experiment was all in order and accurate enough to be able to detect such a difference, this would be headline news for physicists. So, what makes the .005% “close”? Furthermore, why does this small fraction of the graph warrant any extra attention? Each question is answered by taking into consideration the actual numbers that v commonly takes, and juxtaposing this with the empirical value of c . The speed of light is around one billion km/hr. What makes this significant is not the number in the abstract, but instead the fact that the number is so large when we compare it to any reasonable earthbound velocity that is typically encountered. Unless doing astrophysics or working on a GPS satellite, engineers and scientists alike rarely encounter speeds so high. This is why we speak of fractions of c only in very narrow, physics-oriented discussions. So first, then, we observe that the beginning 1% of the graph is, for most human endeavors, the frequently encountered portion. Second, there is an issue of how accurately we are able to detect differences in speeds that we commonly observe. To be able to distinguish a difference of a nanosecond, even in an experimental setting, is in most instances quite difficult. Thus a mathematician examining the nature of the two curves might be blind to the fact that for the overwhelming majority of the velocities *the difference in predicted values cannot be detected experimentally*. This is what makes the curves “close” in the low-speed regime. The two curves, devoid of context and examined in an austere mathematical framework, cannot be said to be close to one another or not. Were one to remark of the two that they possessed the property of “limiting to one another as $v \rightarrow 0$ ”, a nonplussed mathematician would rightly answer “so what?”.

Once we consider the appropriate empirical context, the graph of the two curves informs us about physical circumstances as they relate to historical context. Again we find that there is an epistemic payout, so long as in addition to the mathematics we are sure to include a point about the scientific history surrounding the result. Here we will find the *explanans* to be both the mathematics of the reduction and these less-obvious empirical details. Once they are made clear,

this alternate yet still forceful epistemic goal of “explaining the successes of NM” will be made manifest.

Consider the implicit knowledge that, until the end of the 19th century, there were few experiments possible that could have presented a challenge to Newton’s momentum equation. The technology simply wasn’t available. Certainly there might have been a way to detect a discrepancy given reasonable cleverness, yet there wasn’t even any incentive to go looking for such an experiment. Physicists at the time didn’t have any reason to think that CM would miss the mark when tested at (what were then) extraordinary experimental velocities. Add in the further easily-acknowledged experimental and explanatory success of CM until the 20th century and we can begin to understand why such a result is interesting to scientists. The two curves in the graph tell a story of *why* CM was successful and unchallenged for so long. It informs us about why it continues to be taught in schools and why it is still used by (most) engineers. It even helps to give an explanation as to why Sir Arthur Eddington needed to leave the comfort of Cambridge to go to a small island west of Africa: devising an experiment that would provide a critical test of relativity theory is actually quite difficult.

Again I would like to make notice of the importance of these extra-mathematical contributions. None of these historical factors apply to equations like [3.4], and so we can see why the fact that [3.4] reduces₂ to [3.2] is of no interest. The reason why [3.4] seems ridiculous as a possible representation for momentum is that it may easily be observed to be preposterous in a classroom experiment. Equation [3.4] will be within .005% of [3.3] when $v = 30 \text{ m/s}$. We could figure out that it was empirically flawed after doing experiments with a wristwatch and a baseball. We would not consider [3.4] “close” to [3.3], because they will give experimentally different results at speeds that we readily observe and can measure. [3.4], as a possible measure of momentum, is garbage. Only by attending to the empirical context are we able to get full payout from our multifaceted explanation of why CM was successful, and for so long a time. Up until now we considered examples of preposterous “mature” theories that by [3.1] were nevertheless considered reductions, in an attempt to demonstrate how [3.1] alone would admit too much. Now I will show how the requirements of [3.1] may leave out instances that we would want to claim are reductions, yet by [3.1] would be inadmissible. The last section looked to an objection inspired by Nickles that tried to allow performing limits only on “mature” theories. The response claimed that all theories of momentum would have the property of reducing₂ to

[3.3], as they have the same terminal point and are continuous. To find an objection to this reasoning, it would suffice to envision a scenario that purposed a reasonable, “established” theory that still failed to limit to [3.3]. To do so, imagine an alternate future where, rather than SR, Einstein and those involved in the physics community had come up with the following theory:

$$\mathbf{p} = \mathbf{m}_0 \mathbf{v} / \sqrt{(1 - \mathbf{v}^2 / \mathbf{c}^2)} \pm \boldsymbol{\sigma} \quad [3.5]$$

So long as $\boldsymbol{\sigma} \neq \mathbf{0}$, as $\mathbf{v} \rightarrow \mathbf{0}$, [3.5] $\neq$ [3.3]. To justify how such an equation could be arrived at in the first place, imagine that for quantum reasons – or perhaps even for extra-scientific reasons – physicists believed that all objects were always in motion, perhaps revolving or oscillating very slowly. Here every object would possess some baseline momentum relative to some observer. Let $\boldsymbol{\sigma}$ ’s bounds be small enough that it is very difficult to determine experimentally. Now we have a successor momentum theory that does not have the property of limiting to its predecessor, CM. Although I had previously assumed that a momentum equation that was not $\mathbf{0}$ in the limit $\mathbf{v} \rightarrow \mathbf{0}$ would be outlandish⁶¹, I hope this description makes [3.5] minimally plausible.

I think that cases like [3.5] are very important because they fail to meet the reduction₂ criteria as described by [3.1], but nevertheless *should* be included. [3.5] does not approach [3.3] as $\mathbf{v}$ approaches $\mathbf{0}$, and [3.1] requires that it does. When evaluating theories, we must consider the empirical context that those theories represent. For our example, [3.5] doesn’t limit to [3.3], it limits to all of the continuous curves that terminate in $\boldsymbol{\sigma}$. Without knowing the size of $\boldsymbol{\sigma}$ or what it is meant to represent, we are left in the dark. For a mathematician, the difference represented by $\boldsymbol{\sigma}$ – for purposes of limiting – is immense. The result would be the same if the magnitude of $\boldsymbol{\sigma}$ was **50** times $\mathbf{c}$. Recall that $\boldsymbol{\sigma}$ could be below the bounds of reliable experimental precision; thus it would contingently be the case that scientists would be unable to detect the difference between the two theories. Experimentally, they would be just as content with [3.5] as they would be with [3.2] for the purposes of claiming a reduction₂. For example, were we to have an epistemic goal of a superseding theory being able to explain features of [3.3], then inasmuch as [3.2] did so,

⁶¹ It may appear that by making [3.5] as plausible theory, I have invalidated the argumentation I employed in §3.3. I do not think there is a worry, for in §3.3 I merely needed to assure that the concept of “limiting two functions to one another” admits many ridiculous curves. In this argument I presumed that all functions that were plausible went to $\mathbf{0}$ as $\mathbf{v}$ did. That there exist cases such as [3.5] that lack this property does not threaten the problem of the still large number of curves that are ridiculous measures of momentum yet still limit to $\mathbf{0}$ as $\mathbf{v} \rightarrow \mathbf{0}$.

[3.5] would do so as well. [3.5] is “just as good” as [3.2] when relating to [3.3], in as much as it can accomplish all the goals that a successional relation could ask of it. Considering the goals of “explaining the success of β ” and “transferring confidence to α ” – it certainly seems like this curve thus far does fine at this task.

From an attenuated mathematical perspective, [3.5] does not limit to [3.3] any more than any other ludicrous equation does. Thus if we are to assume that [3.1] describes a reduction₂, there can be no talk of reducing₂, and the discussion ends. However, an astute physicist looking at [3.3] and [3.5] would likely wonder: “how big is σ ?” For if σ is small enough, suddenly the difference may not matter empirically. So we would miss out on a result that does what we would like a reduction to do, yet it isn’t a “physicist’s reduction” due to the restrictions imposed by the limiting requirement. This is why I think that a definition such as [3.1] is inadequate to describe all cases. We have seen the conditions required of a reduction become less restrictive in successive accounts: from demanding an exact agreement between α and β (the models in §1.1) to an analogy between approximate theories (Schaffner and the New Wave in §2.1). This example provides a similar tale of caution for those who would seek to demand every reduction follow [3.1]: there are cases where the mathematical limit fails to capture what it means for one theory to be “close” to another.

How “close” can two theories be? *How* small? What do “close” or “small” even mean in these situations? Our only guide here is the physical representations of the variables and how those values are obtained by practitioners in the real world. So either a momentum equation is continuous and has $\mathbf{p}=\mathbf{0}$ when $\mathbf{v}=\mathbf{0}$, making it automatically reduce₂ to [3.3], or it is discontinuous or has a non-zero momentum when the velocity is zero. My argument is that in the former case there will be curves that shouldn’t be said to reduce to [3.3] but nevertheless do reduce₂ to [3.3], and that in the latter case there are curves that should reduce to [3.3] but do not reduce₂ to [3.3]. For these reasons I think that the “physicist’s reduction” cannot be represented solely by [3.1], as Batterman had suggested; ironically, we find that limiting is in some cases not demanding enough, while in other cases limiting is too demanding. Thus we can conclude that it was good Nickles allowed for “approximations of many kinds” to be employed when making a reduction₂ instead of exclusively focusing on limits.

§3.3.2 Transference of Confidence

One of Nickles's observations was that "it was an important confirmation of [SR] to show that it yielded CM in the correct limit" (1973, 185). I agree. However I think that we must be careful when trying to understand what this confirmation consists in. Above I have argued that the "limit" relation is trivial when applied to *v in abstractum*. But can it be used to transfer confidence in a new theory? Epistemologically, this is an issue about justification: will the demonstration of a limiting relation with the prior theory provide reason to believe in the truth of the future theory? I would first like to make a general argument to show that, in *any* case where there is a limiting relation, we will never be guaranteed of transference of confirmation.

Take two functions, **F** and **G**, over a set of variables, $x_1, x_2, \dots, x_n$, and grant that for some values assigned to the x_i , $F(x_1, x_2, \dots, x_n) \neq G(x_1, x_2, \dots, x_n)$. Let it be that for a non-empty range **R**₁, **F** and **G** disagree, and for a range **R**₂, **F** and **G** agree⁶². It may well be that there are additional characteristics that relate **F** and **G**, such as one limiting to another. The important feature is not merely *that* they coincide (exactly or "nearly") in some cases, but instead that *these* are the cases which matter. **R**₂ needs to include a range of values that are important tests for the theories, values that have justificatory weight. To make a non-mathematical analogy, if I were to be confident in letting a friend teach my logic class, evidence could not merely be that, for some range of answers to questions, we both gave the same answer. We might agree about a large number of answers to banal questions such as: "where is the classroom located?" or "how many letters does 'logic' have in it?" However, I would have concern if my friend could not give reliable answers that were even close to correct concerning *the questions that mattered*, such as "what is a binary connective?" or "is the λ -calculus Turing-complete?" Similar reasoning applies to the ranges **R**₁ and **R**₂. For any **R**₂, without knowing which values are important, that there are agreements tells us very little. This is because it might well be that the agreements occur on points which are experimentally easily-satisfied, experimentally inaccessible, theoretically pathological, etc. Likewise the disagreements that constitute the range **R**₁ need not be cause for concern, under the same contextual criterion. There might be many disagreements, yet the values

⁶² If **R**₁ was empty, the functions would be equal for all sections for which they were each defined. Even in cases of pristine Nagelian reductions, there must be *some* difference between the two theories, for them to be considered different theories.

on which they disagreed were again deviant: they form a set which is not experimentally relevant, are theoretically uninteresting, or that the disagreement in each case is so miniscule that it is not of concern. Here we can imagine a case where, although $\mathbf{F}$ and $\mathbf{G}$ differed for some $\mathbf{x}_i$'s, the difference was experimentally undetectable, such as a difference in position that is a millionth of a percent of the Planck length. So, in the abstract, without any knowledge of *which* points matter, we must remain silent on whether or not there is to be any transference of confidence provided by a reduction (or whatever we might seek to call a comparison of two equations). Instead, as has been the lesson throughout §3.2 and §3.3, we must remain attentive to the context – to the actual physical values to which the points correspond empirically.

Now that we have understood that we must also consider the empirical implications of the equations, we can still try to analyze why Nickles would see the SR-CM limit as an “important confirmation”. I believe that we may reinterpret this locution to be merely a claim about empirical adequacy. SR needs to be successful empirically to become the preferred theory. Its predecessor, CM, had 200 years of successful predictions and confirmations, albeit in a limited velocity range. One quick way of piggybacking off these successes is for SR to assert that it did “just as well” as CM did in this domain, or, “about as well as we could tell”. This is again just a rather colloquial way of asserting that “SR limits to CM for low velocities”, or “in the $0 < \mathbf{v} < 0.01\mathbf{c}$ range, the difference between [3.2] and [3.3] was experimentally undetectable until the early 20th century”. Here the “reduction” merely disguises a transitive inheritance of empirical success: the CM momentum equation was empirically successful over a domain, the SR and CM momentum equations are close enough that we cannot empirically distinguish the two over this domain, thus the SR momentum is empirically successful over the domain.

Had there been no CM momentum equation, but instead a pile of accrued data, this would have sufficed as a “confirmation” in much of the same way that the CM limit did. In this respect CM serves as a proxy for the data that confirmed it. However it does not merely represent base data; CM is a *bona fide* scientific theory, and this counts much more than a mere datum-agglomeration. The CM momentum equation is embedded within this structure. For years CM was subject to many “critical tests”, ones that were designed specifically to question CM in varying situations. Scientists worked hard to falsify CM, and failed to do so in the vast majority of attempts⁶³. Thus the CM momentum equation is robust in that it is well-tested, and thus this

⁶³ For instance, the Michelson-Morley experiment was seen as problematic (1887).

robustness is transferred to the SR momentum equation, at least in part, when we make the comparison that is the limit and when we consider that the experimental difference in the lower-velocity range is negligible.

The specifics of our case should make it obvious why we must again introduce empirical context. [3.2] and [3.3] are in exact agreement only when $\mathbf{v}=\mathbf{0}$ – just once! Furthermore the only means of comparison between them with regards to $\mathbf{v}$ is the limit that the two share, a limit I have argued above to be impoverished. But given the discussed context, this is more than enough to establish transference of empirical adequacy. As a result, I agree with Nickles that the reduction provides a confirmation of sorts, so long as we consider “the reduction” to consist in the mathematical comparison of the two curves *in addition to* the relevant empirical contextual factors.

§3.4 The Problem of Limiting $(\mathbf{v}/\mathbf{c})^2$:

Certain philosophers of science have chosen to view the SR/CM reduction from yet another angle: limiting the quantity $(\mathbf{v}/\mathbf{c})^2$. Rohrlich (1988), Rohrlich and Hardin (1983), and Batterman⁶⁴ (1995) all focus on how to limit this parameter, which is dimensionless and thus allows for discussion that might seem to avoid some of the contextual issues we have previously considered⁶⁵.

When speaking of the limit, Rohrlich and Hardin avoid using the notation “ $(\mathbf{v}/\mathbf{c})^2 \rightarrow \mathbf{0}$ ”, and instead refer to the “strong inequality”: $(\mathbf{v}/\mathbf{c})^2 \ll \mathbf{1}$. As we have not encountered this notation previously, one might wonder if it adds anything new to the discussion, mathematically or conceptually. Batterman is explicit about what such notation entails: it is an order of magnitude estimate that means something very particular in perturbation theory. To claim that $\mathbf{f}(\mathbf{x}) \ll \mathbf{g}(\mathbf{x})$, some limit must be performed on $\mathbf{x}$, i.e. as $\mathbf{x} \rightarrow \mathbf{x}_0$. This is understood to mean that:

⁶⁴ Batterman’s focus for the SR/CM result falls on a consideration of the type of limit that is operating, attending to how close the neighborhood of the limited perturbations are from other unperturbed expressions. I will postpone this discussion until §4.1; until then it will suffice to say that our limit is a normal one under Batterman’s criterion.

⁶⁵ Note that nothing has been neglected by not featuring $\mathbf{v}/\mathbf{c}$, as the discussion about $(\mathbf{v}/\mathbf{c})^2$ is sufficient. Operations, such as limiting, performed to the square of a value will result in the same effect as performing the operation on the unsquared value.

$$\lim_{x \rightarrow x_0} f(x)/g(x) = 0. \quad [3.6]$$

This interpretation is far superior to the casual locution “ $f(x)$ is much smaller than $g(x)$ ”, as it is precise and also allows that both functions may have interesting limiting properties. For example, both $f(x)$ and $g(x)$ may limit to 0, but so long as $f(x)$ does so “faster”, the relationship will obtain.

So how are we to interpret this schema for the $(v/c)^2 \ll 1$ case? Neither Batterman nor Rohrlich fill in the details as to how we are to fill in [3.6] for our current case. We have several options on how to demark the x of [3.6]: we may choose v , $1/c$, or $(v/c)^2$. Both Batterman and Rohrlich have chosen to employ $(v/c)^2$, partly for the ease of making a Taylor expansion about a dimensionless quantity. Doing so forces $g(x)$ to be 1, while $f(x)$ will be $(v/c)^2$. Under this interpretation, [3.6] instantiated becomes:

$$\lim_{(v/c)^2 \rightarrow 0} (v/c)^2/1 = 0 \quad [3.7]$$

[3.7] is the full meaning of the expression “ $(v/c)^2 \ll 1$ ”. Here the extrapolation is rather uneventful, as an assertion that “ $(v/c)^2 \ll 1$ ” is tantamount to claiming that we limit $(v/c)^2 \rightarrow 0$ – the same way that we have understood the limit in the last two sections⁶⁷. However it is still important for us to recognize that, by examining the momentum equation by “the strong inequality $(v/c)^2 \ll 1$ ”, Rohrlich and Hardin have not deviated from the main thread of our discussion.

Rohrlich and Hardin are interested in characterizing the limiting process in terms of a “validity limit”:

A validity limit is thus equivalent to a specification of the error made by using β instead of α . Any predictions by β should be multiplied by a factor $1 \pm \delta$ where δ is an order of magnitude estimate of the error made. (Rohrlich and Hardin 1983, 607)

⁶⁶ I take these definitions from Batterman (1995, 174), wherein he cites Bender and Orszag (1978) as his relevant mathematical source.

⁶⁷ Were we to choose v or $1/c$ as our x for the functions f and g , then we would still be left with a constant for the remaining g or f . As an example, one way of choosing the functions is to let $f(v) = v$ and to let $g(v) = c$; here g is still a constant function and we would read the strong inequality as “ v is much smaller than c ”. Suffice to say that were we to choose v or $1/c$, our discussion would continue along the lines of §3.2 or §3.3.

We now have an explicit, dimensionless measure for how two theories compare. Let $\mathbf{P}_\alpha$ represent the “predictions made by α ” and $\mathbf{P}_\beta$ represent the “predictions made by β ”. By doing so we may condense the definition provided by Rohrlich and Hardin in order to represent how these predictions vary in relation to one another over the range of possible values:

$$\mathbf{P}_\alpha = \mathbf{P}_\beta(1 \pm \delta) \quad [3.8]$$

There are times when the validity limit will vary, namely when $\mathbf{P}_\alpha$ and $\mathbf{P}_\beta$ are more complicated functions. Between [3.2] and [3.3], the separation changes with $\mathbf{v}$. Here δ changes as $\mathbf{v}$ changes, so any consideration of δ should be made with this variance in mind. Other cases are easier to describe, such as those for which δ is constant. Unfortunately in most cases the corresponding δ is rather complex.

What is δ when we compare the momentum equations of SR and CM? The predictions of either theory, concerning momentum, will be made solely on the extant physical conditions, and the momentum equations of each theory. The relevant empirical inputs required by the SR momentum equation and the CM momentum equation will be the same for each case, as we will operate with a fixed $\mathbf{c}$. Thus, for [3.8], $\mathbf{P}_\alpha$ can be substituted as $\mathbf{p}_\alpha$, the momentum equation of SR that is given by [3.2]. Similarly we will use $\mathbf{p}_\beta$, the momentum equation of CM at [3.3], in the place of $\mathbf{P}_\beta$. Before we can read off an equation like [3.8], we must first finesse the two momentum equations. One common and easily-generalizable procedure to make a limiting process clear is to characterize a function by a Taylor (or Taylor-like) expansion, representing $\mathbf{p}_\alpha$, or portions of $\mathbf{p}_\alpha$, as a series of powers of variables. After this is done, the inner-workings of approximations and other operations such as limiting become much more transparent. First notice that:

$$\mathbf{p}_\alpha = \mathbf{p}_\beta(1 - \mathbf{v}^2/\mathbf{c}^2)^{-1/2}$$

we may provide a Taylor expansion (about $\mathbf{0}$) of the right-hand side of the equation so that we arrive at:

$$\mathbf{p}_\alpha = \mathbf{p}_\beta(1 + (1/2)(\mathbf{v}/\mathbf{c})^2 + (3/8)(\mathbf{v}/\mathbf{c})^4 + (5/16)(\mathbf{v}/\mathbf{c})^6 + \dots) \quad [3.9]$$

Notice that this expansion makes it very easy to see that, by limiting $(\mathbf{v}/\mathbf{c})^2$ to $\mathbf{0}$, $\mathbf{p}_\alpha = \mathbf{p}_\beta$. As [3.9] takes the form of [3.8], we arrive at the validity limit between the momentum equation of CM and the momentum equation of SR. Specifically:

$$\delta = (1/2)(\mathbf{v}/\mathbf{c})^2 + (3/8)(\mathbf{v}/\mathbf{c})^4 + (5/16)(\mathbf{v}/\mathbf{c})^6 + \dots \quad [3.10]$$

Reminding ourselves that $\mathbf{c}$ is a constant, we find that δ is a function of $\mathbf{v}$, as expected. Hoping to prevent pathological cases from fitting their schema, Rohrlich and Hardin restrict their analysis to *mature* theories. The authors list four conditions for “maturity”, which I read as being individually necessary and jointly sufficient:

(i) Mathematical Structure

Theories or theory-parts must be represented mathematically so as to generate quantitative predictions that are sufficiently resolute. This in some ways seems a bit restrictive, as it would discount elements of theories from many of the non-physical sciences, such as the function of “fitness” in evolutionary theory. However we can tolerate the requirement because the analysis Rohrlich and Hardin intend only functions reliably for theories that possess such a structure. Also note that this need not imply axiomatization, for the authors note this requirement is rarely achieved and also inessential.

(ii) Empirical Support

The theory must be well-corroborated over known regimes, but additionally needs to provide novel predictions under which it may be tested. In the (brief) words of the authors, the theory must make “predictions beyond the data base on which it was originally founded” (Rohrlich & Hardin 1983, p.604).

(iii) Horizontal Coherence

There must be an agreement “horizontally”⁶⁸: other branches of science must corroborate the claims made by the theory, and the theory must be minimally consistent with descriptions and characterizations made by other well-regarded theories in other areas. The authors avoid any specific explication of how we are to make these inter-level comparisons.

As an example of how the third condition might be realized, the authors cite the “convergence of results from radio carbon dating, derived from nuclear theory, with the results from independent dating techniques in geology and archaeology” (Rohrlich and Hardin 1983, 604). One problem with examples such as these is how they rely heavily on a strong notion of coherence for science. Were one to consider a patchwork of different scientific laws⁶⁹ – laws that need not fit together nicely into an integrated pyramid – then there need not be strong agreement across theories of differing disciplines. Each theory or group of theories might work quite well within a given domain, and might do quite poorly in predicting or cohering with other theories outside of this domain.

Even granting the authors some strong realist/fundamentalist assumptions about science, it turns out that the example they have chosen is itself contentious. Radiocarbon dating, since its advent, has often erroneously calibrated dates due to a natural variation of carbon-14 in the atmosphere. Thus it has required correction from other areas of science, for example by using tree rings and more recently, sediment core readings from lakes (Ramsey et al. 2012). It has been at odds with dates put forth from archaeologists and climate scientists, and typically the uncorrected dates will underestimate the age of bones and other organic material by several thousand years on samples that are older than 10,000 years (Callaway 2012). In this regard, we find that carbon dating *defers to* the results from other areas of science, rather than *converges with* them.

(iv) Vertical Coherence

⁶⁸ The directionality chosen for Rohrlich and Hardin places the different scientific branches of inter-level reductions from left-to-right on a horizontal axis, while the vertical axis, from top-to-bottom, describes successional reductions from newest theories to oldest theories, respectively. Unfortunately, this has reversed the picture provided by the “scientific pyramid” of reduction depicted in this dissertation.

⁶⁹ I will elaborate on each of these positions in §5.2.

The fourth condition is the most restrictive, as it requires that a theory should agree, in some way, with the successes of the theory/theories that came before it. Rohrlich and Harding describe that for CM and SR, this would require:

... SR to be asymptotically coherent with CM in the following sense: there exists a suitable limiting process by which the domain of validity of α is restricted to that of β ; this limiting process must lead to results which are consistent with β ; in optimal cases it may even reproduce the key equations of β . (Rohrlich and Hardin 1983, 605)

How one is to construe a “suitable” limiting process is crucial, a discussion that we have pursued throughout §3.3. This issue aside, sometimes we are to understand “vertical coherence” as a limiting relationship, yet this need not happen in all cases. Generally, limiting may be understood as just one way to go about restricting a domain of validity. For example, when seeking to compare a more sophisticated Newtonian equation of ballistic motion to a simplistic Galilean equation, we would likely “ignore the fact that the earth is rotating”, so as to mitigate Coriolis effects. Here we have made a move to approximate, yet no limiting is involved.

Rohrlich and Hardin are quick to preempt a potential problem if this fourth condition had no further qualifications: if it were required that every α be coherent to a *mature* β , this would necessitate that, among other things, a given β must have been coherent with its predecessors. This generates a mild regress problem: how could the “first theory” ever be considered mature? To skirt this issue, the authors qualify that (iv) only apply in cases where β is itself mature. This allows for a “first mature theory” to emerge unhindered by its relationship to predecessors. It also prevents a worry that any wild scientific theory that was coherent with α over a domain could work, as a mature theory must meet (ii) by being empirically supported.

Up until this point, we have been concerned with the status granted to a current theory, such as SR, by its relation to past theories, like CM. We may say that SR is or isn’t mature based on how well it interfaces with past mature theories. As a general project, I think that the authors have done an excellent job showing that limiting need not exhaust all means of comparing two theories⁷⁰. As we saw in §3.3.1, limiting in the strict mathematical sense requires a great deal – in some cases far more than is necessary when comparing successor theories. Just as Nickles was smart to allow for other approximation methods, Rohrlich and Hardin have cleverly followed suit.

⁷⁰ Although later when speaking on the subject of reduction, Rohrlich falls back to requiring limiting (1988, 304)

Rohrlich and Hardin and Batterman have each decided to examine the limiting of the SR momentum equation from the vantage point of $(\mathbf{v}/\mathbf{c})^2$. Here they used a strong inequality, $(\mathbf{v}/\mathbf{c})^2 \ll 1$, to characterize how $(\mathbf{v}/\mathbf{c})^2$ is constrained. This section demonstrated that doing so is equivalent to imposing the limit that $(\mathbf{v}/\mathbf{c})^2 \rightarrow 0$. Additionally, by expanding the SR momentum equation as a Taylor series, it could be written in the form $\mathbf{p}_a = \mathbf{p}_b(1 \pm \delta)$, where δ is the measure of the error in $\mathbf{p}_b$ from $\mathbf{p}_a$. This measure conveniently allows a succinct expression of how an equation of the succeeded CM differs from the corresponding equation of its successor, SR. Rohrlich and Hardin see this as an important property of a “mature” theory.

§3.4.1 Establishing Past Theories and Applicability

Although interesting, discussing the SR-CM limiting case in terms of $(\mathbf{v}/\mathbf{c})^2$ is somewhat tangential to the purposes that Rohrlich and Hardin would put to a validity limit. The paper was written in the early 1980’s, when the authors were trying to insert themselves into the debate about realism and antirealism in science. One of the worries addressed by their paper was how a realist picture of science could be maintained despite the progression of science featuring numerous past theories that have been superseded, overturned, and otherwise discarded. Rohrlich and Hardin are replying more specifically to Laudan’s *A Confutation of Convergent Realism* (1981). To do this, Rohrlich and Hardin claim that a more detailed look at the history of science will vindicate those of a realist persuasion; specifically they seek to show how certain “established” theories will illustrate a gradual refinement in predictive power as theories progress⁷¹.

Imagine a situation in which there are several theories in contention for the title of “best theory”. Regardless of which one is chosen, each theory has its time in the scientific spotlight, and as science develops, all of them are eventually discarded in favor of a new theory. As a fine example of this, recall the nascent atomic era when several models of the atom were in competition to describe atomic structure. There was a “plum pudding” model championed by J. J. Thomson, which involved electron-like “corpuscles” embedded in a larger positively-charged body (1904). Also competing at the time was a “Saturnian” model developed by Hantaro

⁷¹ There are notable similarities between this project and the structuralist model of theory reduction examined in §1.1.3, as well as the noted “progress” of science that will be discussed in §5.2.

Nagaoka (Hentschel 2009). Here a positively-charged center “planet” was encircled by several negatively-charged “rings”. Although partially developed in reaction to falsifying experiments concerning Thomson’s model, we could perhaps include the Rutherford-Bohr model of the atom, featuring a small positively-charged central nucleus ensconced by electron-inhabited shells (Bohr 1913). Each of these models sought to provide a description of atomic structure, and ultimately each was judged to be erroneous. Thus the opponent of realism could claim that there is “no scientific ‘progress’”, very little coherence from theory to theory, and little that can be said regarding the “approximate truth” of such theories⁷². The maneuver envisioned by Rohrlich and Hardin was to see which of the contemporaneous theories could be considered “mature”, i.e., to see which satisfied (i)-(iv) above. One of these theories would qualify, as clause (ii) requires that it be empirically viable. Once this mature theory was succeeded by a later, more-qualified mature theory, the authors would deem the succeeded theory “established”. Subsequently, we could point to the many established theories as ones that showcase the advance of science, and thereby provide a clear picture of the process the opponent of realism had seen as being otherwise opaque.

Rohrlich and Hardin claim that “when a mature β receives validity limits from an α and thus is reconfirmed within these limits by α , we call it an *established* theory” (1983, p.607, emphasis in original). Just as we saw for “mature” theories, the “established” status is a technical term⁷³, as opposed to the colloquial meaning of “‘recognized and well entrenched since long ago’ which is clearly not intended here” (Rohrlich & Hardin 1983, footnote 5). Notice that because establishment is a status granted to β , and not α , the most current scientific theories cannot be considered established, only at best mature. A fine motivation for ascertaining which past mature theories should be deemed established was to highlight those that played an important role in the history and development of science.

Once a theory is succeeded, the “vertical coherence” it exhibits with its successor is often grounded by a limiting relationship. This limiting relationship allows for a succinct recognition of the validity limit, and thus is reason to recognize a succeeded β as established. As Rohrlich

⁷² The concept of “approximate truth” is important to Rohrlich and Hardin, as they see as a viable measure for characterizing how the predictive power of theories can improve. Admittedly difficult to formulate clearly, “approximate truth” functions quite differently from our traditional received views of “truth”. For example, Arthur Fine shows that we must be careful when conjoining approximately true propositions (1996, 121).

⁷³ As per §3.2.1, “mature” and “established” were both terms employed by Nickles. Here they have a different intention.

considers succession by a more “approximately true” theory to be quite likely (1988, 300–301), there are times when the current mature theory may be known to have flaws, while there is not any mature theory that has yet been devised in replacement. Rohrlich cites “Newton's gravitation theory before general relativity” as a nice historical case (1988, 301). As a result, the status of “established” does little more than distinguish past mature theories from the most-current mature theory.

The SR-CM case follows just as we would expect. Take it for granted that CM and SR are both mature theories. The momentum equation of SR ($\mathbf{p}_\alpha$) limits to the momentum equation of CM ($\mathbf{p}_\beta$) and, as in [3.8], $\mathbf{p}_\beta$ can be said to be within $1 \pm \delta$ of $\mathbf{p}_\alpha$; [3.10] shows us that δ is a function of $(\mathbf{v}/c)^2$. Through the recognition of this relationship between the prior CM momentum equation and the succeeding SR momentum equation – in addition to similar relationships that obtain between the other equations of these two theories that I have not discussed – we are to consider CM an established theory.

An established theory is intended to tell us about a theory's place in the trajectory of science. It is supposed to highlight certain past theories from the history of science as being “close” to the correct theory, showing that science becomes more empirically adequate with each iteration. By following the chain of established theories, we may recover a picture of scientific furtherance in which each theory improves upon the one before it, within the range of nested δ 's.

The roles played by “mature” and “established” theories allow us to present another goal for an intertheoretic comparison. By looking at successional reductions we can see what these relations reveal about the past, succeeded theories. Doing so can tell a story about the history of science, and show us that not all discarded theories are equivalently so. Sometimes multiple theories are in contention, and in the end none are deemed sufficient as a new theory comes in that best explains them all. To say that all of these past theories are “wrong” is to move too quickly. Each is “wrong”, in that they fail to get the correct results in some cases, but all are “close” – and some of these are closer than others. To distinguish which theories should be looked to as the significant successes of the era, we must compare these theories with those that come after them. Only then can we see how well they fit into the march of scientific progress. Thus a reduction that makes explicit the connections and differences between α and β can explain and elucidate the evolution of scientific theories.

Rohrlich and Hardin intended, by looking at a chain of established theories that terminate in a currently-accepted mature theory, to be able to argue for realism on the grounds of the “approximate truth” of past theories. I am ultimately not interested in this project, as it is tangential to my own and has been adequately answered by others (Fine 1996, chapters 7, 8). Rohrlich and Hardin actually disavow any involvement with the reduction project⁷⁴, stating that they “are not concerned here with the question whether a reduction of β to α has been carried out or even can be carried out” (1983, 607). However the concept of employing a dimensionless δ to describe the process of limiting α to β is very valuable, as the next section will demonstrate.

§3.4.2 Recasting Prior Successes

§3.4 showed that the expression $(v/c)^2 \ll 1$ claimed nothing more than “limit $(v/c)^2 \rightarrow 0$, $(v/c)^2/1 = 0$ ”. Notice that, as in §3.2, if $c \rightarrow \infty$, then $(v/c)^2 \rightarrow 0$. Similarly for the work of §3.3, if $v \rightarrow 0$, then $(v/c)^2 \rightarrow 0$. Thus the assumptions of either section will lead to the claim that $(v/c)^2 \ll 1$. In this way, starting from a position of limiting c or v will still result in all of the consequences that usher forth from a claim that $(v/c)^2 \ll 1$, or equivalently that $(v/c)^2 \rightarrow 0$. Much of the benefit provided by employing $(v/c)^2$ is that the Taylor expansion in [3.10] generates an explicit factor δ by which p_α and p_β differ. But before considering these merits, what would it mean, physically, for $(v/c)^2 \rightarrow 0$?

$(v/c)^2$ is dimensionless, and it thus may seem that any limiting of $(v/c)^2$ requires (or perhaps admits) no physical interpretation. Earlier I claimed that we saw physical and historical context enter in when I had to justify what it would mean for $c \rightarrow \infty$, or for $v \rightarrow 0$. Since we are now limiting $(v/c)^2 \rightarrow 0$, how would the magnitude of c matter empirically? The value of c in the world may not matter for understanding the limit, yet there is still a way of understanding what $(v/c)^2 \rightarrow 0$ could mean.

If $(v/c)^2 \rightarrow 0$, this means that the term becomes so negligible so as not to matter for our calculations. $(v/c)^2$ becomes “infinitesimal” or “very small” and has “little to no influence on the equation in which it is employed” – mathematically, $(v/c)^2 \ll 1$. For this to happen, v must be

⁷⁴ Rohrlich later makes clear his thoughts on reduction (Rohrlich 1988). Batterman best summarizes Rohrlich’s position on the SR-CM reduction, using Nickles’s terminology, as: “[CM] reduces₁ to SR because the *mathematical framework* of SR reduces₂ to the mathematical framework of [CM]” (Batterman 2002a, 79).

small relative to $\mathbf{c}$, meaning that $\mathbf{c}$ dominates over $\mathbf{v}$ in such a way as to make their ratio negligible. In what physical situations would this obtain? Precisely ones where we are using velocities that are very small relative to the speed of light. Now the value of $\mathbf{c}$ is relevant, for otherwise we would not be able to discern the range of values $\mathbf{v}$ would take in these cases. Acknowledging that $\mathbf{c} \approx 10^9$ km/hr, and that for the observations and experimental capabilities of the 18th century, $\mathbf{v} < 10^4$ km/hr, then it does appear that $(\mathbf{v}/\mathbf{c})^2 \ll 1$. So the actual values of $\mathbf{c}$ and $\mathbf{v}$ seem to factor in after all.

But this is not where the story ends. None of this yet requires limiting, and it does little to showcase what work is being done by δ . The claim is that $(\mathbf{v}/\mathbf{c})^2 \ll 1$, when representing earthbound, pre-modern velocities, cannot be adequately analyzed without reference to the equation that the limit is to be taken in. To show this, take two hypothetical theories: a succeeded theory $\mathbf{v}$ and its successor theory μ . Let the momentum equation for $\mathbf{v}$ be $\mathbf{p}_\mathbf{v}$, and let the momentum equation for μ be $\mathbf{p}_\mu$. Now presume that the relationship between the momentum measures can be characterized by:

$$\mathbf{p}_\mu = (1+10^{16}(\mathbf{v}/\mathbf{c})^2) \mathbf{p}_\mathbf{v} \quad [3.11]$$

This relationship is meant to be quite extraordinary. For example, when $\mathbf{c} = 10^9$ km/hr and $\mathbf{v} = 10^3$ km/hr, $\delta = 10^4$ – an error that is rather noticeable, as it is of a greater order of magnitude than the original $\mathbf{v}$ itself. However if we were to limit $(\mathbf{v}/\mathbf{c})^2 \rightarrow 0$, $\mathbf{p}_\mathbf{v} = \mathbf{p}_\mu$. And this shows the reason why we would be loath to claim that $(\mathbf{v}/\mathbf{c})^2 \ll 1$ for [3.11]. “Limiting ϵ to 0”, or “ ϵ being very much smaller than 1” cannot be said to have physical significance unless we are given the equations to which ϵ applies, as well as the contexts in which these equations are to be employed, and under what range of ϵ . [3.11] provides an instance where the existence of a limit is vacuous without the physical context.

The above construction does nothing more than recast the old points of §3.2.1 and §3.3.1 with new metal. But observe the ease with which the goal of explaining the success of CM may be demonstrated. Returning to our familiar α of SR and β of CM:

- (1) When $\mathbf{v} < 10^4$ km/hr and $\mathbf{c} = 10^9$ km/hr for [11], $\delta < 10^{-10}$.

(2) Experimental capabilities before the 19th century had problems creating speeds greater than 10^4 km/hr, or detecting an error of 10^{-10} in observed speeds.

(C) CM went unchallenged experimentally before the 19th century.

We find that the validity limit directly relates to the experimental accuracy in a way that may succinctly explain the success of the past theory. Additionally, it shows us why CM is not a “dead” science. It still provides an acceptable method of calculating momentum, so long as we recognize that the results will only be valid within δ .

§3.5 Conclusion

One of the first authors to discuss the role of limiting in reduction was Nickles. §3.1 details the distinction he makes between a reduction₁, a “philosopher’s reduction”, and a reduction₂, a “physicist’s reduction”. The interest falls on the latter notion, to examine exactly how performing approximations, such as limiting, operate in a reduction. To this effect the chapter provides an extended discussion of one of the most referenced cases of a physicist’s reduction that involves limiting: the SR-CM equations for momentum.

§3.2 discusses the case from the perspective of limiting c to ∞ . Nickles was wary of limiting physical constants, but it was argued that this is permitted, so long as there is a physical understanding of the limit that justifies the process. The goal associated with the limit, elaborated in §3.2.1, is that the reduction explains the past successes of CM: why it was successful for so long, and why SR took so long to rise to prominence. This is atypical of epistemically-minded reductions, as the *explanans* is the reduction and the *explanandum* are contextual details relating to each theory.

The process of limiting v to 0 is the focus of §3.3. The section argues for the vacuity of any reduction that contains a condition that the two equations limit to one another. Firstly it shows that there are many preposterous equations that nevertheless limit to one another by virtue of being continuous and also possessing the same endpoint. The section also provides an example of two equations that seem prime candidates for a reduction, but cannot be limited to

one another. To salvage the role of limiting, the context of the limit must be taken into consideration. The physical values taken along the limit, as well as the historical circumstances of either theory, are instead the relevant factors that make the limit important for the reduction. Once these factors have been included, the goal of explaining features of CM and SR's development and history may be realized. Additionally, §3.3.2 shows how the past successes of CM may be used to transfer confidence to SR over the appropriate regimes.

Lastly §3.4 examines the limiting of $(v/c)^2$ to 0. Rohrlich and Harden's usage of "strong inequalities" is ultimately shown to not introduce anything beyond what the other authors had done when limiting values. The authors employ a "validity limit" to signify how much the predictions of the prior theory may differ from the more current theory. This provides a succinct way to demark the domain of applicability for the past theories, as well as providing a metric to judge the predictions of past and current theories. Finally, the same goal of explaining the context and success of CM is also realized.

Chapter 4 - Intra-level Relations: Wave Optics

and Ray Optics

My main motivation for studying intertheoretic reduction is not so much to try to make general claims about the nature of reduction as it is in understanding the particular and peculiar connections and correspondences between certain pairs of theories.

Robert Batterman (1995, 172)

§4.0 Introduction:

Our past focus has been on reductions and models that might seek to describe them. This chapter presents a case that has not been seen as a *reduction*, but instead as a theory *relation*, while discussing the work Batterman has done concerning limiting, theory relations, and theory reductions. In §4.1 we see how Batterman distinguishes between two types of limits, and notice what this implies for intra-level comparisons. §4.2 showcases a discussion of how to explain the universality of certain characteristics of a rainbow using wave optics and ray optics. Here Batterman claims that both theories are needed to give an accurate account of the empirical phenomena. We consider an objection to Batterman in §4.3, where Belot claims that, contrary to Batterman's assertion, only wave optics is necessary for doing successful rainbow physics. Batterman's response is the focus of §4.4. He agrees with Gordon Belot that (in a certain sense) ray optics is contained in wave optics, yet claims at the same time that this does little to provide an explanation of universality. Lastly in §4.5 we look at an objection by Redhead claiming that Batterman has reified the ray-theoretic role in the discussion, when instead it could be regarded as a mere mathematical device. To conclude, §4.5.1 shows how an intra-level comparison can exemplify typical scientific behavior, by providing a scientific explanation of the empirical.

§4.1 Batterman and Limiting:

In §3.1, we saw that Batterman considered a reduction₂, the “physicist’s reduction”, to be best represented by the following relationship:

$$\text{Lim}_{\varepsilon \rightarrow 0} \alpha(\varepsilon) = \beta \quad [4.1]$$

Here we limit a parameter (or several) that are present in α so that we may arrive at β . Batterman is critical about the universality of the reduction₂ model, as he believes it is not always representative of a typical scientist’s activities. To make this point, he provides several detailed examples from physics where one theory is limited and then compared to another, yet the relationship between the two cannot be described simply as one of equality. Instead, by his detailed investigations Batterman shows the mathematical activity constituting this comparison can allow for novel “borderland physics” to emerge.

According to Batterman, one of the easily identifiable deficiencies of the reduction₂ model is that it fails to identify which type of limit is being employed. For example, notice that the limiting employed with c in §3.2 and v in §3.3 can be said to differ in their physical justifications, and additionally, the limit performed on $(v/c)^2$ in §3.4 is different mathematically, as it is dimensionless. Batterman wants to make a further distinction between “normal limits” and “singular limits”, because he feels that they operate very differently when being used to compare theories. When limiting a variable x of a function $f(x)$ to a value k , if $f(k)$ is “fundamentally different in character” than all of the values $f(x)$ takes as it approaches k , call this a *singular limit*⁷⁵. If we find that there is little difference in the “character” of $f(k)$ and all the $f(x)$ values in the neighborhood about k , call this limit a *normal limit*. Batterman contests that when singular limits are involved, a limiting relationship should not warrant a reduction – neither a reduction₁ nor a reduction₂. Indeed, upon providing the details of examples in optics and mechanics that involve singular limits, Batterman refers to each as “intertheoretic relations” as opposed to “inter-/intra-level reductions”. Batterman has championed this change in to avoid some of the baggage that normally accompanies conceptions of reduction.

⁷⁵ Batterman takes his definition from (Bender and Orszag 1978, 324).

Batterman wants to claim that when there is a singular limit, “*no reduction* of any sort can obtain between the theories” (2002a, 5). Yet notice that the GR limit presented in §2.1 *is* singular. As λ approaches 0, every value taken will correspond to a GR spacetime. The light cones open so as to allow a greater possible range for the motion of an object, but not in a way that changes the character of spacetime to be something that would not be a solution to the GR constraints. Conversely at the limit, the result is *not* a GR spacetime but a classical one. In other words, due to the light cone being completely flattened, such a spacetime could not be a solution to Einstein’s equation. The spacetime of geometrized NM allows there to be a notion of simultaneous events with absolute distances between them; such notions are nonsensical for GR spacetimes. Michael Redhead also makes note of this fact, stating that taking the limit $1/c \rightarrow 0$ of GR “is an example of a singular limit in the relationship between two theories” (2004, 528). So according to Batterman the result should not be considered a reduction. I have however shown that the GR/NM case fits exceedingly well to the New Wave reduction model. Hooker refers to the result as a reduction, while Batterman would disagree. Malament and Weatherall have remained silent on the issue. So who is correct? Is it a reduction or not?

Perhaps Batterman is worried about whether a singular limit could possibly support confidence in a transfer of ontological or epistemic commitment. So in this way he is hesitant to use “reduction”, because the term often carries strong ontological and epistemic connotations. But hasn’t the New Wave of §2.1.1 recognized that all reductions need not invoke such connotations, or at least with such intensity? Recall that advocates of this model believe that the strength of the *analogy* dictates whether ontological retention or replacement is appropriate. But the act of taking a limit, singular or otherwise, occurs before the deduction of α^* , as it is part of C , the conditions that are imposed upon α . Batterman is concerned with the distance between α and α^* , as a singular limit would create an analogy that was too weak. No doubt, we arrive at α^* by a derivation, but it is a derivation that occurs *after* α has been suitably transformed by the imposition of a limit that is singular. Batterman’s worry is that such an imposition is too much.

For Nickles, a reduction₂ allows the application of a limit or “approximation of some kind” to α . The ontological and epistemic goals that sometimes surround a reduction are *not* essential to a reduction₂, by Nickles’s account. So Batterman could claim that a singular limit provided a reduction₂, if his worry were that a singular limit would be taken to accomplish these goals. But Batterman also believes that a reduction₂ also will not fit. He is worried that even with

a successional reduction, we should be cautious in allowing all types of limiting. Just as Nickles (wisely) disallowed that *every* type of approximation could be employed in a reduction₂, Batterman would like to make a similar distinction for limiting. Singular limits presumably transform α too much: when the smoke clears after a singular limit, whatever is left to compare – via mapping or derivation or analogy – to β is simply no longer comparable to the original α .

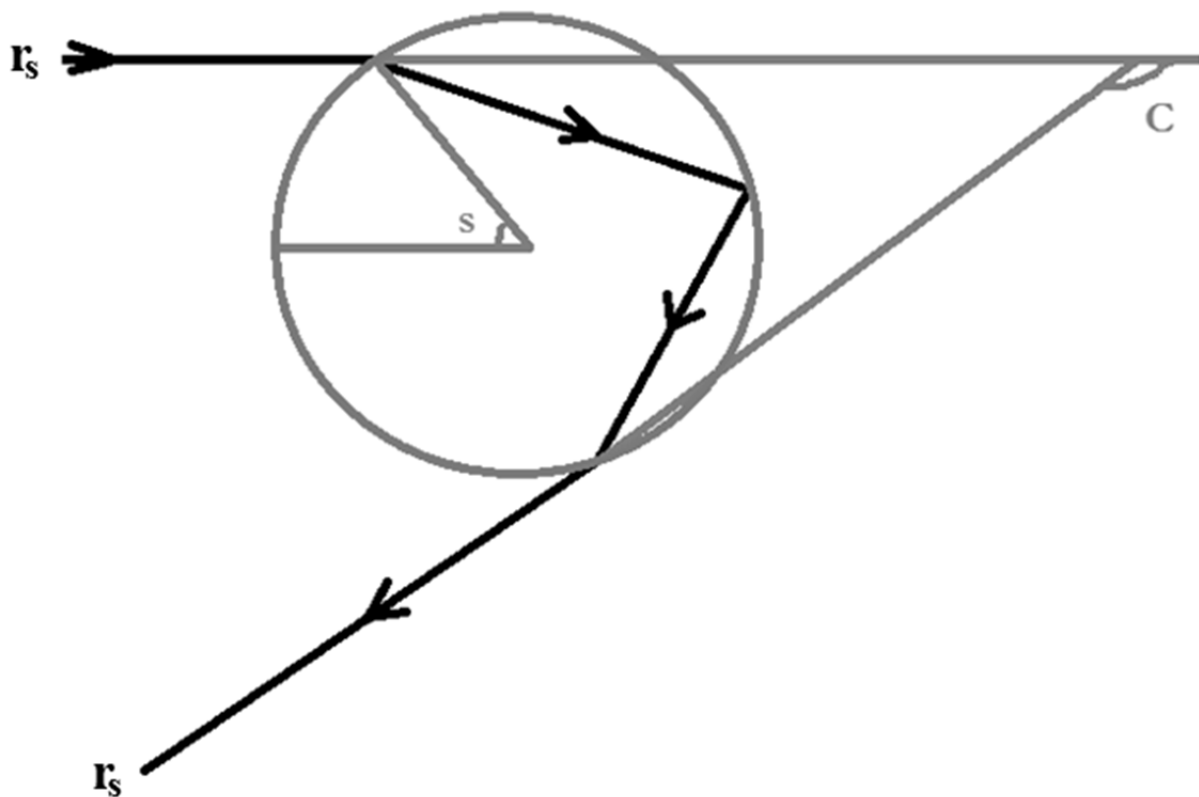
§4.2 Catastrophe Optics

Batterman examines how both wave and ray/geometrical optics are employed in examining rainbows – visual phenomena that are created when light strikes rain droplets. Historically ray optics preceded wave optics; indeed it was Rene Descartes who first showed that by using beams of light and refraction laws, one could explain many of the key properties of rainbows (Descartes 2001, 332–345). Since that time geometrical optics has been developed, and in its modern formulation it is sometimes referred to the “shortwave limit” of wave optics, where the frequency $\lambda = 0$ or the wavelength $k = \infty$. Speaking somewhat metaphorically, taking either limit has the effect of “stretching” the wave out so that it becomes a straight line, or a ray. Batterman’s thesis is a simple one: to explain certain features of the rainbow, neither wave nor ray optics alone is sufficient. Instead we must carefully navigate the asymptotic region between the two theories to arrive at the desired results.

The observational qualities of rainbows are quite interesting. Early in the development of optics, these phenomena provided robust explanatory challenges for early theories in the developing science of light. For example, geometrical optics can describe precisely what angle a rainbow will occur at to an observer. Prior to geometrical optics, an explanation of why rainbows can be seen only at certain angles was unavailable. Although our discussion will apply to a few specific features of the rainbow, I will first give a general introduction to rainbow optics⁷⁶. Rainbows occur naturally only when it is raining, as they are a result of the interplay of the sunlight and myriad raindrops that are present in the sky. When light travels from one medium to another, it will typically both reflect, bouncing off the boundary at an angle, and refract, crossing the boundary after being deflected at an (often different) angle. To begin to describe the process

⁷⁶ When providing the scientific details throughout this section, I rely heavily on Batterman’s own presentation in (2002a, 88).

of how a rainbow is formed, imagine a light ray approaching a spherical water droplet. Some of the light will reflect off the surface and never penetrate the droplet. However the light that does enter will refract, and then potentially reflect several times inside the droplet before exiting. We will only concern ourselves with the light that makes one inner reflection before refracting out of the droplet. If light reflects more than once, it is responsible for creating atypical rainbows, rainbows that are usually quite faint and have different qualities, such as reversed bow colorings. Such multiple reflections will not contribute to our image of the original rainbow, and the secondary rainbows are typically quite faint when observed by the unaided eye.



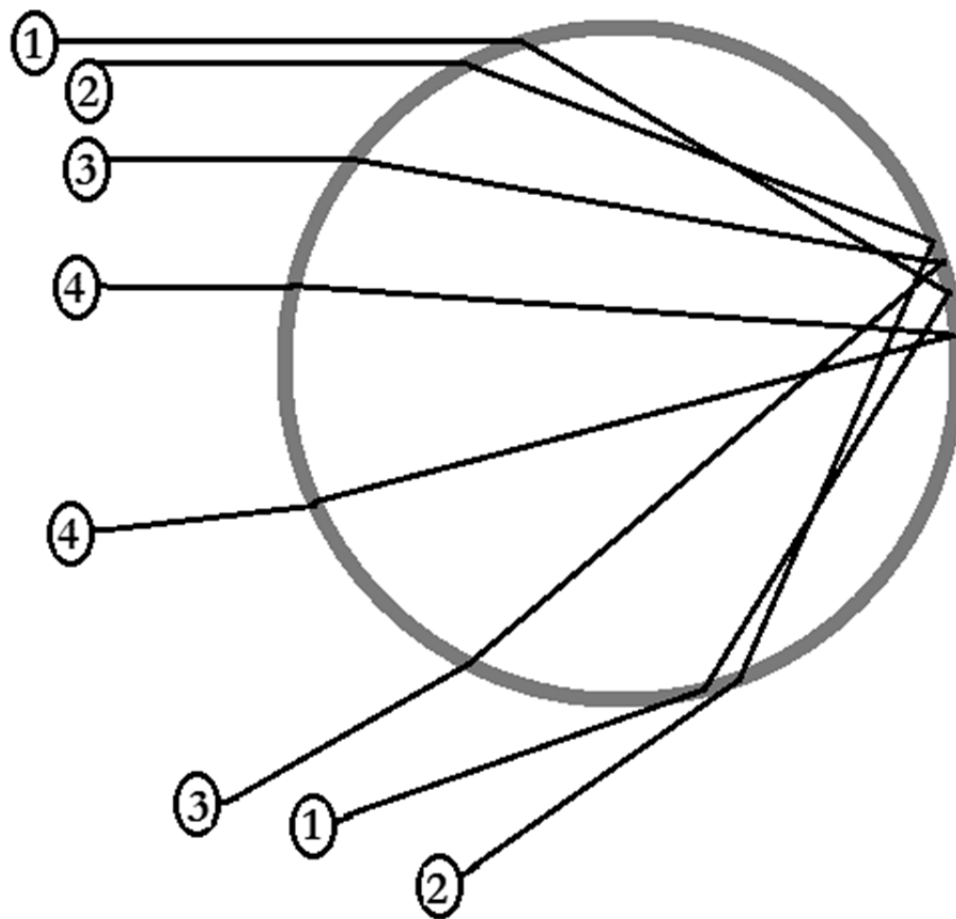
{Figure 4.1}

We may label the incoming rays by the angle s that they make to the center of the raindrop as they initially refract as r_s , as is shown by {Figure 4.1}. As such, we shall focus on: $0^\circ < s < 90^\circ$. We will be interested what happens to the angle of each r_s as they exit the drop; measure this by the angle C that each ray makes to a line parallel to that made by each ray prior to entering the

drop. One may employ Snell's law and basic geometry to arrive at $C(s)$, which will tell us which angle each ray exits at. Using **1.33** as a measure of the refractive index of water, we arrive at:

$$C(s) = 2s - 4\sin^{-1}\left(\frac{1}{1.33}\sin(s)\right) + 180 \quad (\text{John A. Adam 2002, 244}) \quad [4.2]$$

Several different configurations of C and s are shown in {Figure 4.2}. Importantly, C is minimal at $C_0 \approx 137^\circ$. Call the ray that achieves this minimum the “rainbow ray”. On {Figure 4.2} the rainbow ray is labeled as **2**.



{Figure 4.2}

Geometrical optics provides a nice description of why rainbows occur. Around the rainbow ray, light rays “stack up”. This convergence of rays results in a rapid increase in

intensity: as we approach C_0 from 180° , the intensity blows up. This is because of all the possible incoming rays, r_s , a large number of them will exit at an angle very close to C_0 . In relation to {Figure 2}, this means that there will be many rays that are very near to the rainbow ray that is 2. This will make the rainbow appear bright to the observer.

Although the ray theory provides an excellent account of why rainbows occur, and at which angles, it cannot provide an adequate description of what the intensity should be at or near the rainbow ray. This convergence of light rays is called a “caustic”. Caustics are singularities of ray theory: “the caustic is a line on which the intensity of light is, strictly speaking, infinite” (Batterman 2002a, 88). Here is a problem for ray theory, as it cannot account for what the intensity of light is at or near the caustic.



{Figure 4.3}⁷⁷

⁷⁷ Photo credit to Andrew Dunn - <http://www.andrewdunnphoto.com/>; I have changed the image to black and white.

An additional feature that cannot be explained by ray optics is the presence of “supernumerary bows”. On the inside arc of a rainbow, just below the last violet band, alternating bright and dark bands appear. As the bands become more distant from the rainbow, their width decreases, causing them to slowly taper off. In ideal conditions, the light and dark bands of the supernumerary bows are visible to the naked eye, as shown by the contrast-enhanced photo that is {Figure 4.3}. Indeed, an early success of the wave theory was to provide an explanation of why supernumerary bows occur (Young 1804). By being attentive to interference patterns that occur when light is treated as a wave, George Bidell Airy was able, in 1838, to develop a wave-theoretic equation that determines the intensity of light at and near the caustic (1838). In [4.3] the equation is written as a function of y , where y is a function of C (and thus a function of s)⁷⁸:

$$Ai(y) = \frac{1}{\pi} \int_0^{\infty} \cos\left(\frac{t^3}{3} + yt\right) dt \quad [4.3]$$

To find the intensity, we take the modulus squared of [4.3], $|Ai(y)|^2$, after substituting the additional factors relating to the physical particulars of the drops. A graph of this intensity is shown in {Figure 4}. This function, it turns out, is successful on many levels. First, it predicts a finite and empirically-corroborated intensity for the caustic. It also predicts that for positive values of y there is a non-zero intensity that tapers off, an effect that is visible in rainbows as a slight glow that resides on the outside of the bands. The peak of the intensity on [4.2] is in fact a little below the rainbow ray as portrayed by {Figure 4.1}, not centered directly as the ray theory postulated. This discrepancy with the ray prediction is also empirically vindicated: we observe the angles of the bands slightly shifted. Most importantly, the oscillations that taper off in the $-y$ direction indicate that there will be bows other than the first ones. These are precisely the supernumerary bows.

⁷⁸ Specifically, $y = 4^{1/3} k^{2/3} a^2 (n^2 - 1)^{1/2} (4 - n^2)^{-2/3} (C - C_0)$, where a is the radius of the drop, k is the frequency of light, and n is the refractive index of water, and C_0 is the rainbow ray angle $\approx 137^\circ$ (Jackson 1999, 34).

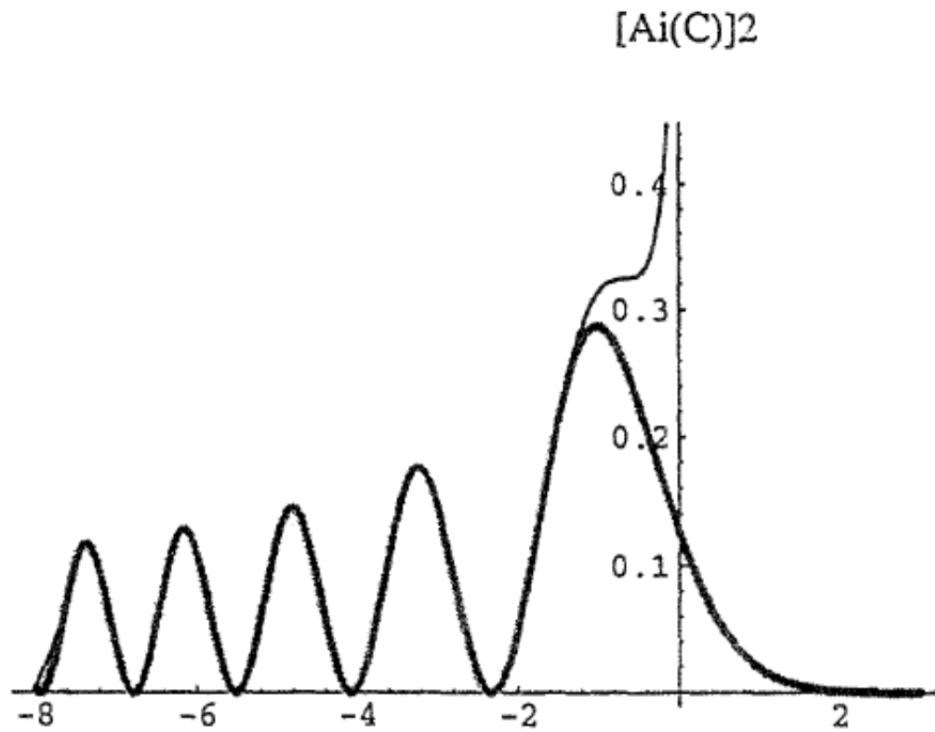


Fig. 5a. Square of the Airy function.

{Figure 4.4}⁷⁹

It may look like any conversation concerning comparison is over, because the wave theory provides a superior description of the phenomena in question to the one provided by the ray theory. But we are not done yet. Here is the complication: were we to vary the *shape* of the drop slightly, we would find that *all* of our above reasoning would need to be redone. A change in the shape of the drop would throw off all of the angles— equation [4.2] would no longer be accurate. This would result in a new C_0 , which you recall is required for us to employ [4.3]. Yet, the Airy integral, as applied to the rainbow problem, *assumes* a constant radius for the drop, i.e., it presumes that we are working with a sphere. We know empirically that raindrops are not perfect spheres (Beard and Chuang 1987) (Pruppacher and Klett 2010, sec. 10.3.2). The hope, of course, was that the model created with a spherical drop could be fit reliably to real-world rain droplets. However, once we recognize that a different droplet will result in different light-reflection patterns, there is no guarantee that these will play out in ways similar to the spherical

⁷⁹ This graph was taken from (Batterman 1995, 183). Along with Batterman, I am less interested in the values of the axes and more with the shape of the function. As such I have left the axes unlabeled just as Batterman did.

case. And if we are not confident that our model is empirically applicable, this serves in turn to undermine our explanation of real-world rainbows, as well as to prompt questions about how we can model variations of droplet shape.

If the goal is to provide an explanation of the *universality* of fringe spacings, Batterman notices that we cannot simply perturb some parameter of our model to generate a local variant of [4.2] or [4.3] (Batterman 2002a, 90). Instead some other entirely different approach is required. Batterman takes his cue from Michael Berry and uses resources from the mathematics used to describe caustics: catastrophe theory. Catastrophe theory focuses on finding critical points of functions; when used to model caustics it effectively treats light as rays, for there is no reference to any wave properties. Instead catastrophe theory merely considers the nature of curves that appear at the confluences of line intersections. Catastrophe theory allows us to characterize a caustic by its codimension $\mathbf{K}$, where $\mathbf{K} = (\text{dimensionality of the control space}) - (\text{dimensionality of the singularity})$. The result relevant to our case is that for $\mathbf{K} \leq 7$, caustics are stable over diffeomorphism (Batterman 2002a, 91). It can be shown that any spheroid droplet (that isn't too distorted) will give rise to a caustic of the sort described by [4.2]. There will be a singularity in each case, each with a corresponding picture subtly like the one given by {Figure 2}. The stability of elementary caustics means that any smooth perturbation of a droplet shape will result in a caustic that has the *same essential structure*. For our case, $\mathbf{K} = (1) - (0) = 1$. Thus every caustic will be a member of the class described by the $\mathbf{K}=1$ caustic, the *fold caustic*.

The structure of the elementary caustics may be represented by polynomials. For the fold caustic, we use the following:

$$\frac{1}{3} \mathbf{s}^{1/2} + \mathbf{C}\mathbf{s} \quad [4.4]$$

When we employ this polynomial to model a wave function, we arrive at the following:

$$\Psi_{\text{fold}}(\mathbf{C}) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} e^{-i(\frac{\mathbf{s}^3}{3} + \mathbf{C}\mathbf{s})} d\mathbf{s} \quad [4.5]$$

Close to (or at) the caustic, we are able to describe intensity by [4.5], by taking the absolute square of the wave function, $|\Psi(\mathbf{C})|^2$. The intensity of the real part of the wave function will yield

the same intensity as that given by the Airy integral, i.e. $|\Psi(\mathbf{C})|^2$, will equal the intensity of [4.3], $|\mathbf{Ai}(\mathbf{C})|^2$. This is to say, in addition to showing the universality of fold caustics as droplet-size varies, employing mathematics from catastrophe theory allows us to capture some of the same results from the wave-theoretic description of the rainbow. The connection between ray and wave optics is quite nuanced, as results such as these demonstrate.

Thus we have an explanation of why the intensity patterns are universal. Batterman goes on to show how catastrophe theory, when combined with wave-theoretic premises, predicts how the fringe patterns will scale with intensity. Each pattern will differ as the frequency of light increases, and we know that when $\mathbf{k} \rightarrow \infty$, intensity likewise goes to infinity. Looking at [4.3], we know that $\mathbf{y}$ is a function of $\mathbf{k}$. Also, it can be shown that as $\mathbf{k} \rightarrow \infty$, $[[4.3]]^2 \rightarrow \infty$. But this description only worked for the case of the perfectly spherical raindrop. To arrive at a general formula, start with the wave-theoretic equation⁸⁰ for the rainbow wavefront:

$$\Psi(\mathbf{C}) = \left(\frac{\mathbf{k}}{2\pi i}\right)^{1/2} \int \mu(\mathbf{C}) e^{ik\varphi(\mathbf{C})} d\mathbf{s} \quad [4.6]$$

In this equation, substitute the general form of the fold catastrophe, given by [4.4]. Next make a simple linear transform, $\mathbf{k}\mathbf{r}^3 = \mathbf{s}^3$. Once this has been done, we arrive at:

$$\Psi(\mathbf{C}) = \mathbf{k}^{1/6} \mu \int a(\mathbf{C}) e^{ik\varphi(\mathbf{C})} d\mathbf{s} \quad [4.7]$$

Notice that this equation captures many factors, such as how the patterns of intensity will scale with increasing $\mathbf{k}$. [4.7] tells us that the intensity will increase as $|\mathbf{k}^{1/6}|^2 = \mathbf{k}^{1/3}$. Importantly, this scaling occurs for all rainbows⁸¹, as we used the ray-theoretic catastrophe theory to produce the result.

Ray theory alone could not adequately describe certain phenomena, such as the observed intensity of the light coming from a rainbow. To do so we needed to take into account wave-theoretic properties of light. However wave theory was unable to describe how perturbations on

⁸⁰ Here $\mathbf{C}$ and $\mathbf{s}$ are the control parameter and state variable, just as we have defined them by {Figure 4.5} above, $\mu(\mathbf{C})$ represents the amplitude, $\varphi(\mathbf{C})$ the “optical distance” function, and $\mathbf{k}$ is of course the frequency.

⁸¹ If there are drastic changes to shape of the droplet, such as making it a cube, naturally rainbows will not properly form. Furthermore it is known that other significant deformations to the shape result in aberrant rainbows, such as the “dewbow” or “fogbow” (John A. Adam 2002, 237).

the shape of droplets would nevertheless result in a similar result for the supernumerary bows' intensity pattern. To do so we needed to employ a result from catastrophe theory, and as such, we were forced to treat the light as rays. Additionally, equations from both theories were needed to arrive at a description of how the intensity scales with light frequency. All of these considerations are what led Batterman to conclude that, when describing how caustics are employed to analyze and explain rainbows, it is not a simple picture of a successor theory providing a pristine treatment of the cases where the succeeded theory has failed. Instead, "in the asymptotic domain characterized by the limit $\lambda \rightarrow 0$ or $k \rightarrow \infty$, *elements from both theories seem to be intimately and inexorably intertwined*" (Batterman 2002, 94, emphasis in original). Without examining at the details one might easily assume the immediate primacy of the successor theory. Instead by carefully attending to cases such as this we are able to show a complicated picture involving both theories, "new" and "old".

§4.3 Belot's *Ab Initio* Objection

Central to Batterman's thesis was the claim that, to arrive at an explanation of the universality of rainbow phenomena, and model the intensity of caustics, both ray and wave optics are necessary. In the paper "Whose Devil? Which Details?", Belot attempts to arrive at the desired results by employing *only* wave optics in his modeling. Because he proceeds by giving an *ab initio* treatment, Belot asserts that, mathematically, geometrical optics is contained in the equations of wave optics – any result garnered from geometrical equations could be achieved by some succession of operations performed upon the wave equations. Logically, therefore, it could not be that "two theories" have been employed, because every consequence of ray optics is also a consequence of wave optics. And granting that there are many results about the rainbow that ray optics cannot account for, such as the existence of supernumerary bows, one can see the inclination to award wave optics status as the superior, more "fundamental" theory. This complements the many accounts of theory comparison that take reduced or succeeded theories to be special cases of the reducing, succeeding theories⁸².

⁸² Each of the three models of reduction from §1.1 effectively had this as a consequence. Also there is a similar quality in Schaffner's analog theory β^* of §2.1, the New Wave's α^* from §2.1.1, as well as Rohrlich and Hardin's picture of succession in §3.4 both come to mind.

Recall that Batterman saw invocation of ray optics as unavoidable due to catastrophe theory's ability to account for the appearance of diffeomorphic caustics for a large range of spheroidal droplets. This is what gave confidence in the universality of the rainbow explanation. Belot admits that there are not "any overarching theorems" available to wave optics to show the generality of the result and that likewise, even by constructing wave-theoretic models of several spheroidal cases, a comparison of these solutions to those of the perfect sphere case would almost certainly do little to bolster our confidence in any compelling general conclusions. Belot's main push, however, is that the important perturbations performed on the relevant integral in the high-frequency limit – those which ultimately show universality – still do not essentially involve a reference to ray optics or the theory of caustics. To make this point explicit, Belot shows how one could construe the entire operation as a speculative exercise put to a pure mathematician. Here one has merely to observe the effect of constraints and boundary conditions of a manifold parameterized by the relevant differential equation. By doing so, the problem is reduced to a mathematics problem concerning the behavior of a specific configuration space at various extrema⁸³ and a mathematician undertaking this project could arrive at a number of considerations important to the result.

Handed a certain set of initial and boundary conditions [the mathematician] has shown that in the high frequency limit, Airy's integral provides a good approximation for the behavior of the quantity called "the intensity of light" for a certain region of space; she has shown, perhaps via numerical integration, that the pattern of "the intensity of light" matches the curves illustrated in the textbooks for "the rainbow"; finally, she has shown that certain qualitative aspects of this pattern of "the intensity of light" are invariant under a certain family of perturbations of the mathematical problem that she has been set. So far the analyst has only a mathematical understanding of the problem. (Belot 2005, 25)

For these reasons Belot is confident in claiming that ray optics is unnecessary. Furthermore this should come as no surprise, he believes, as "rays and caustics are implicit in the apparatus of the wave theory" (Belot 2005, 19).

The last observation is a point that Batterman himself was well aware of when he initially wrote on the subject. He acknowledges that the results in question can be "predicted from [wave optics]", but it is the meaning of this locution that he questions:

⁸³ I have omitted the technical discussion of how this proceeds. To see these details, I would refer the reader to (Belot 2005, 19–24). Suffice it to say that nothing seems incorrect or controversial concerning the mathematics; Batterman tacitly agrees, as he takes no issues with these details in his response to Belot's article (Batterman 2005).

So “predictable from fundamental theory” is somewhat ambiguous. In one sense, the solutions are *contained in the fundamental wave equation*. This is the sense in which asymptotic analysis enables one to find mathematical representations of those solutions. On the other hand, the understanding of those mathematical representations requires reference to structures foreign to the fundamental theory. In this sense, they are *unpredictable from fundamental theory*. (Batterman 2002a, 96)

So who is correct here? Both parties admit that the pure mathematician is able to “find” the universality of the behavior of the integral under perturbations. But what is it that must be added to “understand” them? Batterman thinks that in gaining this understanding, ray optics still has a significant role. This is precisely the issue that he and Belot fight over.

In a trivial sense, the understanding of *any* physical theory – *qua* explaining and describing the physical world – must make reference to considerations beyond those of pure mathematics. Constituents of the theory correspond to empirical entities, and minimally to make note of this correspondence we will be forced to step outside of the mathematician’s ivory tower. However I do not think that Batterman or Belot are contesting this point. Instead it is a question of how much we are to grant when constructing the bridge towards a legitimate empirical explanation. Belot recognizes the charge that he may be “[mistaking] bloodless formalities for genuine physical understanding” (2005, 3), and as such is explicit about what he thinks must be additionally granted to the mathematician:

(i) to impart to [the mathematician] the standard sense of “the intensity of light”; (ii) to explain why the given initial and boundary conditions correspond to a situation in which a cloud of spherical water droplets is illuminated by white light; and (iii) to explain why the perturbations studied correspond to changes in the shape of drops. (Belot 2005, 25)

Belot takes each of the above concessions not in any way to signify a committal to ray theory. Mathematics aside, this is the crux of Belot’s argument, so we will deal with each of the three clauses in detail.

The first clause, talking about how light intensity works, how it varies, and how it scales, appears quite unassuming. This can be done solely in the context of wave optics, and indeed some problems about intensity are only obfuscated when we attempt to explain them in ray-theoretic terms. The second clause seems to be little more than an assertion of a correlation of mathematics to the empirical world. We are simply telling our mathematician that she has

unknowingly constructed a model of a rainbow scenario. There are undoubtedly many points in this conversation that are non-trivial, but each seems to involve details about how the world works and not invoke any machinery beholden to geometrical optics. The mathematician might ask “how does refraction work?” or “by considering the wavefront as a continuous entity, have we done an injustice to the viability of our model?” Good questions, but it is difficult to see where wave optics could fail to do as well as ray optics in providing answers.

The third clause is a bit more complicated. The overarching goal is to arrive at an explanation of the universality of the rainbow physics model when applied to real-world rainbows. In examining just the details of the spherical droplet case, we are left with two options if we are to claim that we have explained how rainbows – in general – work: either we must claim that all rain droplets are spherical, or give an argument that will reliably extend confidence to spheroidal cases. Just as Belot and Batterman both attempt, we must explain why spheroidal models still behave much like the spherical case. It would be quite unreasonable to claim that we had to provide all such models, for there are simply too many. This problem should not be terribly surprising to anyone who has sought to construct covering, universal explanations in a D-N fashion. As a matter of fact Belot recognizes that a similar issue is present for **(ii)**, for “one cannot consider infinitely many sets of possible initial and boundary conditions” (Belot 2005, 18). Thus any explanation must take substance via an argument for generality, not as a mere enumeration. By showing that the range of wavefronts possible due to perturbations of spheroidal droplets correlates with the variations made upon our pure mathematician’s manifold, we should be assured of the coverage of the case. Here again it is difficult to see how **(iii)** will involve recourse to ray optics. The discussion about the matching of perturbations inside an integrand need not defer to “rays” or their ilk.

Consequently, Belot concludes that the reliance of the rainbow explanation on geometrical optics is a fiction. Wave optics alone is sufficient to provide the result. Indeed, a mathematician confronted with the relevant differential equations could arrive at all of the appropriate conclusions via clever asymptotic analysis. Then, by introducing only a few naïve correlations with these equations to the physical world, she would have given a universal explanation of rainbow phenomena. Thus we are left to conclude that Batterman is bedeviled by details, albeit inessential ones.

§4.4 Batterman's Contextual Response

Batterman responds to Belot, taking issue over what theoretical edifices are to be considered “essential” in the treatment of rainbow-associated phenomena (both Belot’s objection and Batterman’s response appear in the same issue of *Philosophy of Science*). Although satisfied with the technical aspects of Belot’s discussion, Batterman nevertheless believes the ray theory is shouldering some of the explanatory labor. Batterman sees geometrical optics as clandestinely operating during the connection of the mathematics to the empirical domain: specifically in the connections (ii) and (iii) attributed to Belot above.

Batterman believes the wave theory to be “explanatorily inadequate”. Wrapped up in this claim is the idea that asymptotic explanation “is really quite different” from other scientific explanations that have been discussed by philosophers of science (Batterman 2005, 155). It may well be that the ray theory is “contained in” wave optics, but, as we have seen, this does little to temper how the explanation is to proceed. Belot’s move was to show how mathematically, once constraints were imposed on the relevant differential equations, the asymptotic analysis may proceed without ray-theoretic obstruction. But Batterman takes issue with how the “constraints” or “boundary conditions” get their beginnings. He claims that “we must examine the physical details of the *boundaries* (the shape, reflective and refractive details of the drops, etc.) in order to set up the *boundary conditions* required for the mathematical solution to the equation” (Batterman 2005, 159). Here the physical aspects dictate how Belot’s pure mathematician is to proceed. Most importantly, these physical aspects – what parameters to vary and how much variance to permit – necessarily involve ray-theoretic understanding.

Some of the boundary conditions were devised by considering what would result in representing a variance in the radius of the droplets. Doing so is necessary to satisfy (iii). But this requires an observation that the envelopes of light are unchanged in their general appearance once they exit these perturbed drops. Any talk of “how the light reflects off the back of the raindrop and refracts as it enters and exits the drop” (Batterman 2005, 159) involves a consideration of ray theory. It could not, at this stage of the problem, be waved away as being described adequately by a more basic wave optics, *because we do not yet know which boundary conditions to impose upon the wave theoretic equations*. Thus, to even get the asymptotic

analysis off the ground, and to understand how and what to vary, we must have some pre-existing notion of how light behaves through mediums – how it refracts, how it reflects, and how it converges. This is precisely the understanding that ray optics facilitates. Thus Batterman argues for the theory-ladenness of the boundary conditions: to know how to set up the constraints, the physical modeling of the scenario involves a ray-theoretic perspective. It is here important to remember that discussions of fundamentalism are often concerned with what is theoretically necessary. Reductionism then follows as one way to realize this goal. Is there one theory that can be used to describe the world, with one subatomic ontology and a few simple laws? Showing that the trajectory of a cannon is guided by microscopic laws is quite a difficult applied problem, even in one specific case. So instead the plan of attack involves reducing the known laws that govern ballistic trajectories of macroscopic semi-rigid bodies to those laws that govern smaller constituents of cannon-stuff. And so on with these laws and objects all the way down to the “fundamental” laws. Such a process will give confidence, supposedly, to claims of our macroscopic laws and descriptions being convenient yet dispensable in principle. Even if such a program is difficult to accomplish for sufficiently general examples, counterexamples to the “fundamental” character of the microscopic laws seem more argumentatively accessible: were we able to show that a given set of macroscopic laws or phenomena couldn’t be a consequence of their purported microscopic correlates, then an inference to dispense with any hopes of fundamentalism seems viable.

We see that both authors are enticed by considerations of what entailment amounts to. When one acknowledges that ray optics is “a consequence of” wave optics, logical relations like the following come to mind: if ϕ entails ψ , and ψ entails χ , then it must be that ϕ entails χ . By similar reasoning, it appears that by conceding that wave optical theory entails ray optical theory, one is thereby committed to claiming that any explanatory work done by ray optics must likewise, in principle, also be able to be done by wave optics. But this is an urge we must resist: notice that logical/mathematical entailment is transitive, yet “entailment” as it relates to explanation is not transitive. Perhaps if our notion of explanation is a Hempelian deductive-nomological one, we might be able to push the transitivity through with some additional work – but recall that one of Batterman’s important theses concerning the rainbow examples was that *asymptotic explanation cannot be adequately captured by the deductive-nomological model* (1997, 407–408, emphasis in original). So especially for this case, the fact that ray theory is

entailed by wave theory does not necessarily provide justification for inheritance of explanatory features.

Another issue central to the debate concerns the pieces that should be considered necessary to the explanation – issues of what is indispensably employed and at what stage. Batterman claims that ray theory is required to motivate principled decisions concerning how to constrain the wave equations. Strictly speaking, ray theory wasn't part of the mathematical analysis that followed once the problem had been delineated, but Batterman insists ray theory was operant during the “setup” stages of the problem. Critical to explaining the universality of certain aspects of rainbow phenomena was showing how varying the droplet shape leaves the “shapes” of the ray-envelopes invariant (Batterman 2005, 159). Doing so does not simply involve fiddling with a simple function that describes the radius and tracking this fiddling throughout several well-defined equations. Instead, we must make a non-trivial jump by inferring that anytime we “perturb the phase function of the Airy integral” (Belot 2005, 140), the corresponding models will be those that correspond to variances on the radius of the droplet. And what gives us confidence in *this* maneuver is precisely what Batterman wants to highlight as being “(ray) theory-laden”.

§4.5 Redhead's Accusation of Reification

When writing his *contra*-Belot article, Batterman also envisions himself as responding to Redhead. This is because Batterman sees Redhead as presenting an objection very similar to Belot's. I do not think this is the case. Here is the relevant passage from Batterman:

I am being accused [by Belot] of improperly reifying the mathematical structures of the superseded emeritus ray theory when I claim that such structures are required for genuine physical understanding. This objection has also been raised by Michael Redhead (2004) in his discussion of my book. (Batterman 2005, 155)

Directly afterwards Batterman quotes Redhead:

What Batterman is effectively doing is to reify this auxiliary mathematics so that the ray structure becomes part of the physical ontology of a new third theory inhabiting what he calls the “no man's land” between the wave and ray theories. But why can't we leave the asymptotic analysis of universality at the level of a purely mathematical exercise? This

would be in line with other developments in theoretical physics where surplus mathematical structure with arguably no physical reference is used to explain or “control” what is going on in a physical theory. Modern gauge theories are an obvious example of this sort of thing. (Redhead 2004, 530)

If we view Belot as making an accusation of reification, it would be on the grounds that the mathematical, physical, and conceptual maneuvers beholden to ray optics have been improperly attributed as being necessary to the treatment of the problem. From this perspective, Batterman may be said to reify in the sense of “asserting as essential”.

I see Redhead as using “reification” in the sense of the “fallacy attributing more robust existence or presence to an idea or a concept”; this is why he refers to the “physical ontology of a new third theory” in the above quotation⁸⁴. The issue on the table is whether Batterman is justified in claiming that geometrical optics, with its caustics and rays, is joined with the phases and interference patterns of wave optics to create a third theory ontologically inhabited by members of each ray and wave optics.

As we saw in §4.4, Batterman contests that the roles played by ray theory are employed when imposing boundary conditions to wave equations. Thus he thinks that they are essential to the explanation, and the charge of reification in the sense attributed to Belot is appropriately engaged by his discussion. However, I think that this is not the case for the sort of reification that worries Redhead. To answer this question we must ask what it means to be a member of a theory’s ontology, for mathematical constructs as well as physical ones.

At this stage of the discussion we may concede, along with Batterman, that certain maneuvers beyond the canonical governing equations of wave optics are required to motivate the asymptotic analysis performed in the explanation of the various rainbow features. We may further admit that such maneuvers were originally obtained by imagining that light moved as a straight line, among other considerations. These features are furthermore historically considered to be within the purview of ray optics. Still, the salient issue for those involved is the status of the claim that, by employing these considerations and techniques, we are thereby committed to

⁸⁴ To dispense with the obvious, it should be quite clear that by employing the locution “physical ontology”, Redhead is not asserting a realist position about what exists in the world. He is instead speaking of the ontology of a *theory*, the things that the theory itself posits to be the relevant objects, forces, etc. For instance, an electron is a constituent of the ontology of particle physics, while a gene is a construct that we refer to in the ontology of genetics. Whether there “really are” electrons or whether protons may be said to be “more real” than genes are matters completely tangential to these attributions.

taking ray optics – *qua* physical theory – to be a portion of the physical theory that governs the phenomena.

There seem to be three options regarding the status of the asymptotic analysis and its ray-like underpinnings. First, we could claim with Redhead that the theory of wave optics is employed in the explanation, along with some mathematical inspiration from ray theory. Next we could claim that there are two full-blooded theories that are each indispensably operant in the explanation: ray optics and wave optics. Lastly, we could instead claim that there is just one theory being used, and it is an amalgam of both ray and wave optics, for it draws at different points from both of the pieces. Batterman takes the third option: claiming that there is a “third distinct theory in this asymptotic domain between the two giants, wave and geometrical optics” (1995, 185). He is hesitant to embrace the second option because several of the equations in our explanation employ “elements from two incompatible limiting theories” (Batterman 1995, 185). It is additionally not the case that in the explanation each theory is applied at a different point, or that each concerns a different “level”. Instead both theories are needed in the *same* explanation: “if we are trying to understand and explain certain features of the rainbow, ... *both theories seem to be operating at the same level*” (Batterman 2002, 116, emphasis in original).

Superficially, the difference between the “two old theories” and “one new theory” perspectives might seem to hang on a consideration of what a theory can be said to be comprised of. Are we to consider “differential calculus” as a constituent of electrodynamics, or should we think of it instead as “math” that merely helps characterize the physical theory? Depending on how one answers this question, it might seem that our answer to Redhead’s objection will likewise be decided. As we shall see, I think that a distinction such as the one drawn between the second and third options is ultimately not one that should bother us, and the reason for this need not fall on our characterization of a “theory” in regards to the mathematics that surround it. I believe that we can leave this latter issue aside and still arrive at an answer to Redhead.

In favor of Redhead’s position, imagine a world where ray optics had never been created. Instead, physicists had arrived at the governing tenets of wave optics by insight that had no ray-optical influences. In this case, when attempting to explain the supernumerary bands of the rainbow – ceding with Batterman – we will allow that these physicists would be stuck. Now let’s presume that one of the physicists had a casual interest in the pure mathematics of asymptotic analysis and caustic stability. Here this physicist serendipitously realizes that the techniques may

be of use to the problem, cleverly applies them, and churns out the explanation as Batterman has provided. This scenario is an interesting inversion of the situation that inspired Batterman's jeremiad against Belot: instead of a pure mathematician tacitly relying on the ray physics, we have a practitioner of wave optics relying on the abstract mathematics of caustics.

The situation is not strictly inverted, as the "setup" of the problem involves thinking of light operating as straight one-dimensional lines. But this still need not be overly troubling. When asked by her colleagues about what motivated her setup of the problem, our mathematically-inclined theorist could merely claim that it was "a different way to think about how to impose constraints". Further fiddling with the equations when limiting the wavelength and the resultant stability of perturbations at the caustic could be attributed to insight gleaned from her decidedly-unphysical caustic analysis. Thus it seems that we could claim an explanation, and it would only involve wave theory in combination with some conceptual and mathematical help from an abstract catastrophe theory. Any insistence that "light was a ray" would thereby appear as fallacious reification, vindicating Redhead's objection.

I would claim that the difference in the status of the involvement of ray optics or ray-like mathematics in the two worlds is merely circumstantial. For Batterman, the case presented in our actual world contains a "new third theory" that exists between "two differing theories". However in the other possible world, the one where there was never a ray theory, Batterman would have to describe the act of adding catastrophe theory to wave theory as being a "significant expansion to wave theory". So significant that the old wave theory, uninformed by the mathematics and modeling insight of catastrophe theory, would be a genuinely different theory. The status is therefore uninteresting, as any difference between the two worlds to be a consequence of the *de dicto* status of "ray optics" and "ray mathematics". The *de re* status of the various theories and how they were related to one another is ultimately unchallenged. *Something* has been added that significantly modifies wave theory, and I conclude that it is not important if we claim that it is "ray theory" or "ray mathematics". The issue instead concerns how we are to consider these rays.

But what is light? A ray or a wave? Considered as a question about that ontological status as light *qua* physical entity, I am unmotivated to answer. As I mentioned at the onset of this section, the more interesting question is: what we are trying to accomplish by asking about the ontological status of light, as it functions in the most successful theories that describe it? But here we have a quandary: sometimes we will treat light as a ray, and sometimes we will treat it as

a wave – which should we say that it is, *qua* theoretical entity? Just as wave and ray optics are “incompatible theories”, at first glance it looks like light must be either a wave or a ray. For the high-frequency limit is a *singular limit*, and as such we cannot say that both theories agree in that limit. One of the appeals of Redhead’s account is that it makes sense of what our best theories say about light: it is a wave, and when we treat it as a ray we do so as a mathematical maneuver to afford an otherwise elusive explanation, nothing more.

In the history of modern physics, this is not the first time such an issue has come up, even when considering light. Should we consider light as a wave, or as a particle? A very fast answer would be: it is both. When pressed to explain the oft-confused wave-particle duality, a physicist would do well to say that our best theories can explain and predict scenarios involving light only when we consider it at times as a particle and at times as a wave – indeed there are even instances where we must do both simultaneously. This has confused many who are then tempted to press on about the status of light, not as it functions in our best theories, but concerning what it “is really like”. Here the philosophically-informed physicist would do best to provide an answer, one that is certainly beyond the scope of our discussion. The important thing to note is that, when confronted by the wave-particle duality, no one is tempted to claim that light-as-particle physics is foundational and that light-as-wave physics presents a mere mathematical artifice, or vice-versa. Instead we are fine admitting that our explanation of cases like the dual-slit experiment is necessarily of a different character than the majority of our macroscopic explanations involving people walking through doors. Likewise Batterman could claim that we rely on two differing conceptions of light (wave and ray) simultaneously. Catastrophe theory will then “of necessity make reference to both ray theoretic and wave theoretic structures in characterizing its ‘ontology’” (Batterman 2002a, 119). We do so because each is necessary for the explanation, and any seeming inconsistency about the ontological posits of this borderland theory should leave us likewise nonplussed.

Redhead charged Batterman with reifying ray optics. Redhead claimed we could consider any consideration of “light as a ray” as mere mathematical aid, giving wave theory exclusive purview over ontology. I argued that nothing important hangs on viewing the theory employed in explaining the rainbow as “a third theory between ray and wave optics” or “a new type of wave optics”. I believe that, either way, we do rely on a model that attributes a ray structure to light, “treating light as a ray”. But I do not see this treatment as being necessarily inferior to the

portions of the explanation that “treat light as a wave”. The only reason why we might hesitate to claim that it does both would presumably issue from concerns of consistency, as it would seem odd to call light “a ray” and “a wave” in the same explanation. However modern physics has given other examples that reject this worry, for we can treat light as either and still be consistent. As we have relied on a treatment of light as a ray while describing the scenario, I am content to claim that our theory has “light-rays” in its ontology. This does not seem like reification, but instead an accurate description of the role played by light-rays in the explanation.

§4.5.1 Intertheory Comparisons as Scientific Activity

Traditional models of reduction champion one theory as “succeeding” another: when being employed to describe phenomena the succeeding theory will in all cases offer a superior treatment to the succeeded theory. One of Batterman’s central points was to show that cases such as the one examined in §4.2 do not fit this schema. After pursuing an accurate description of the rainbow, we produced “neither laws of the reduced nor of the reducing theories” (Batterman 2002a, 95). An examination of the intra-level hinterland between ray and wave optics shows not a quagmire infested with poorly-justified bridge principles, but instead a fertile valley, one where Batterman’s “catastrophe optics” blooms. Something new and significantly different from either theory emerges from the comparison, and stands out as a clear benefit of it. Likewise, when looking at the details of limiting quantum mechanics to classical mechanics, Batterman reveals that there is a lurid nether-realm he deigns “semiclassical mechanics” (Batterman 2002a, 95). Within semiclassical mechanics are resources to begin to unpack the multifarious relationship between the two larger theories, something that neither theory could do adequately on its own.

One of the important results from the intra-level comparison between ray and wave optics is that new, interesting science may emerge that genuinely involves components of each theory. Often this will occur when the boundaries between the two theories are sufficiently complex. Another of Batterman’s claims is that for certain theories, in cases where an asymptotic treatment of each cannot be easily managed (such as when they possess singular limits), a clever amalgamation of the two is sometimes required. Comparing succeeded theories to their

successors has many benefits; as this case of rainbow asymptotics demonstrates, one of these benefits is the creation of new theory that adds to the body of successful science.

Could this then be considered a goal for a reduction? Unsurprisingly, much of the issue hinges around what a reduction amounts to. For the case of §4.2, it appears to be more a question of how we can use *theory* (ray optics, wave optics, or some combination of the two) to explain *observations*. This is very different from the cases in §2 and §3, where the relationship we were concerned with was primarily *theory-theory*. Part of Batterman's message was that the case of §4.2 is not a case-study detailing a reduction, but instead an "intertheoretic relation". One of the lessons seems to be that we should not be too quick to dispense with prior theories, for they may indeed have an important explanatory role to play in contemporary science.

But wouldn't this imply that we are making a terminological error by referring to ray optics as the "succeeded theory"? The point of §4.2 is to show that it in fact was needed, so our habit of referring to the theory as "prior", "past", or "antiquated", should in fact be revised. In some ways it may be apt to view catastrophe optics as a combination of the two theories employed to describe recalcitrant phenomena. Another take is to see catastrophe optics as a successor to either of the two theories. However construed, by limiting two theories we are able to answer scientific problems that we could not otherwise.

What if Belot was correct in his *ab initio* construction of rainbow phenomena from wave optics alone? Perhaps there will be the emergence of an "overarching theorem" which provided an explanation solely from wave theory. Here we can nevertheless concede that Batterman has demonstrated an explanation of supernumerary bands that resources tools from both ray and wave optics, albeit inessential ones. This would provide a different avenue to achieve the *explanandum*, one that surely could be useful in cases where it was didactically beneficial or otherwise more accessible to those without the mathematical resources that Belot needs to draw upon⁸⁵. So the theoretical comparison is still of benefit; Belot would only disagree with Batterman about the uniqueness of the *explanans*.

If Redhead were correct, and we were to view the machinery from the ray theory as mere mathematics, then admittedly we would not have added to the theoretical body of wave optics proper. We might contest that we have extended the capabilities of the antiquated ray theory,

⁸⁵ Admittedly there is not much of a difference in mathematical difficulty in the two constructions, but we can still envision how a student pursuing Belot's construction might find the process overly erudite and detached from the phenomena in question.

perhaps as historical curio or as a theory that could be helpful in an established validity limit of the sort discussed in §3.4. But even this is not essential for Redhead, as he argued that nothing new is added to the discussion by our asymptotic analysis beyond novel mathematical tools. For Redhead the exercise of explaining the universality of rainbow phenomena was fecund, inasmuch as it provided an application of new mathematics. This would be problematic for Batterman’s thesis, because the role of “ray optics” in this case is repressed. Furthermore, it would cease to be a proper inter-level comparison. It is not “inter-level”, for the role of catastrophe theory is now recognized as being purely mathematical, nor is it “a comparison”, as we are merely employing new techniques to provide an explanation of empirical phenomena. However in §4.5 I gave reason to disagree with Redhead, and additionally showed why the classification of the case is likewise uninteresting.

Overall, I think that the case discussed in §4.2 exemplifies what occurs in most scientific activity. There exists a recalcitrant phenomenon, the presence of supernumerary bands, and there is a turn to theory to provide a (covering) explanation. The dominant theory, wave optics, is inadequate at providing this explanation. So theoreticians must do novel work to generate it. To do so they rely on the help of either another physical theory (Batterman’s belief) or additional mathematics (Redhead’s suggestion). Whether this help was essential in principle is what Belot contested. In each case, the work done is merely the furtherance of scientific knowledge, a non-trivial yet typical endeavor. Thus an inter-level relation can provide an explanation, not of progress or of past theories properties, but a *scientific* explanation of the empirical. This goal is different from the ones we observed in §2 and §3, and provides an important example of a case that can involve the limiting of a “succeeded” theory and its “successor”, yet not be an example of a reduction.

§4.6 Conclusion

This chapter provides details of an asymptotic intra-level comparison that was the focus of a debate between Batterman and Belot. The chapter began in §4.1 by investigating the distinction Batterman makes between a “regular limit” and a “singular limit”. Limits of the latter type occur when there is a significant difference between the behavior of the function *as it*

approaches the limit and its behavior *at the limit*. Batterman claims that we should only consider an asymptotic relationship a reduction when the limit involved is regular.

The central case concerned phenomena associated with the rainbow: the universality of fringe-spacing structure and the intensity patterns around caustics. §4.2 summarized Batterman's presentation, in which he sought to show that neither ray optics nor wave optics played a privileged role in the explanations of these phenomena. Instead, Batterman maintained that each theory was indispensable to explaining the universality of fringe spacings. As a result we cannot view the case as one where the most current theory is solely responsible for the explanation; there exist cases in contemporary science where "succeeded" theories still play a significant role.

Belot responds in §4.3 by attempting to do exactly what Batterman said could not be done: explain the universality of the desired rainbow features from the theoretical resources of wave optics alone. Belot believes that a pure mathematician could operate on the un-interpreted equations of wave optics and achieve the same results. Thus Belot believes that any reliance on ray optics is simply a curio of Batterman's explanations, and not essential in any meaningful sense. §4.4 discusses Batterman's response to Belot. Batterman takes issues with the setup of Belot's equations, claiming that ray optics has been relied upon to impose boundary conditions upon these equations, i.e. to determine where and how to constrain the equations. Thus Batterman reasserts the necessity of ray optics in the result.

The last voice in the discussion is Redhead's. §4.5 examined his accusation of Batterman reifying ray optics. Redhead argued that, even if one is to concede the role of catastrophe theory in the discussion, we can still judge its function to be purely mathematical. In this way one need not view the moves made in §4.2 as having been "interpretations of light as a ray", thereby invoking ray-optical theory, but instead as analytic aides to working with the wave-theoretic equations. I disagree, as I claim that there if we are to employ a consideration of light as a ray in a significant way when setting up our problem, it is not illicit to claim that this was a posit of our theories' ontology. We should not necessarily be worried about a theory that treats light as a ray and then later as a wave, as such issues have arisen in modern physics before.

The chapter concludes with §4.5.1, where we examine what goals underlie the rainbow case. Despite involving "succeeded" and "successor" theories, limiting, and non-trivial mathematical moves, the case is ultimately deemed to be one that is typical to science: theory explaining phenomena. Asymptotic, intertheory activity in sciences need not be exclusively

driven by issues of reduction. Instead Batterman's case questions the dispensability of "past" theories when examining complicated empirical cases, such as the one provided by the rainbow.

Chapter 5 - Reductive Models, Ordering, and Scientific Structure

It is not the role of the working scientist to make his fundamental assumptions clear, and it is not reasonable to expect that he will proceed according to rigorous logical rules. Hence, it is customary for the philosopher of science to consider science in an idealized form.

Kemeny and Oppenheim, (1956, 6–7)

§5.0 Introduction:

Reduction is a dyadic relation, typically one whose relata are each theories, parts of theories, or in some cases mere equations. In §5.1, the discussion will examine the formal features of several models of reduction, specifically checking to see whether each is *reflexive*, *antisymmetric*, or *transitive* with respect to the reduction relation. These three relations will allow us to conjecture about ways that theories may be *ordered* relative to one another. The section will investigate: Nagel's model (§5.1.1), Kemeny and Oppenheim's model (§5.1.2), Suppes's model (§5.1.3), Shaffner's model (§5.1.4), and finally the New Wave model (§5.1.5). §5.2 turns to consider how the goals of a reduction relate to these ordering relations. §5.2.1 argues for the vacuity of reflexivity, on the grounds that it has no strong relation to any of the goals a reduction would seek to accomplish. Antisymmetry, discussed in §5.2.2, is crucial for any model that seeks to include notions of successorship or directionality; consequently symmetry is seen as a problematic quality with regards to many concepts of reduction. §5.2.3 considers the importance of transitivity showing progress across time as well as compositionality from the macro to the micro scale. Finally, §5.2.4 explores how these three relations may be employed to order theories in science. This additionally allows for an examination as to what philosophical aims might underlie those who would endorse a given model. Here I also discuss what further implications the different ordering relations have when considering the unity, or disunity, of science.

§5.1 Models of Reduction and Three Ordering

Relations:

Recall that §1.1 examined a methodology for assessing the success of a model of reduction by Niebergall. He maintained that “any attempt to explicate ‘ β is reducible to α ’ can only be acceptable if (i) it is faithful to examples, and (ii) it is a sound representation of general ‘properties’ of reducibility” (Niebergall 2002, 148). Hopefully the last 4 Chapters have provided ample examples to test the various models of reduction. My goal in this section is to examine (ii). I have chosen three properties to assess:

- (a) The reduction relation is *reflexive*: α is always reducible to α .
- (b) The reduction relation is *antisymmetric*: if β is reducible to α and α is reducible to β , then it must be that $\alpha = \beta$.
- (c) The reduction relation is *transitive*: if γ is reducible to β , and β is reducible to α , then γ is reducible to α .

Before I justify why I have chosen these relations, I will mention how they relate to how theories may be ordered.

If I had a set of three theories $S = \{s_1, s_2, s_3\}$, and a binary relation R , there are 9 possible ordered pairs between the three elements of S , there are 2^9 possible ways that R might be defined by extension. Some of these R ’s are nowhere reflexive, in that they do not allow that the theories relate to themselves. $R_1 = \{ \langle s_1, s_2 \rangle, \langle s_2, s_3 \rangle, \langle s_2, s_1 \rangle \}$ provides such a case. Additionally R_1 isn’t (everywhere) antisymmetric, as symmetry holds between s_1 and s_2 . Some R ’s that are antisymmetric also are not transitive, such as $R_2 = \{ \langle s_1, s_2 \rangle, \langle s_2, s_3 \rangle, \langle s_3, s_1 \rangle \}$. But notice that there is no *ordering* of the theories in R_2 , as if R was “taller than” or “younger than”. As the set is finite, R_2 ’s members cannot be ranked as follows: for every member s_i , there is an s_j such that $\langle s_j, s_i \rangle$. For this to happen both antisymmetry and transitivity need to hold, such as in $R_3 = \{ \langle s_1, s_2 \rangle, \langle s_2, s_3 \rangle, \langle s_1, s_3 \rangle \}$ or $R_4 = \{ \langle s_2, s_3 \rangle, \langle s_2, s_1 \rangle, \langle s_3, s_1 \rangle \}$. Since antisymmetry and transitivity

both occur, each of these $\mathbf{R}$'s could be represented in a manner which showcases that ordering: (s_1, s_2, s_3) for $\mathbf{R}_3$ and (s_2, s_3, s_1) for $\mathbf{R}_4$.

The orderings for $\mathbf{R}_3$ and $\mathbf{R}_4$ are additionally asymmetric, defined (again in terms of reduction of theories) as:

(d) The reduction relation is *asymmetric*, i.e., if β is reducible to α then β is not reducible to α .

A relation can be ordered if it has two properties: transitivity and either asymmetry or antisymmetry. I have chosen to focus on antisymmetry rather than asymmetry to allow for the possibility of reflexivity. Asymmetry cannot happen if a set is reflexive, but antisymmetry still allows for an ordering while remaining silent on whether a given $\mathbf{R}$ is reflexive or not. For example, the relation $\mathbf{R}_5 = \{ \langle s_1, s_1 \rangle, \langle s_2, s_2 \rangle, \langle s_3, s_3 \rangle, \langle s_3, s_1 \rangle, \langle s_3, s_2 \rangle, \langle s_2, s_1 \rangle \}$ is reflexive, antisymmetric, and transitive, and is also ordered as (s_2, s_3, s_1) , but $\mathbf{R}_5$ is not asymmetric. Lastly notice that if a relation is asymmetric, then it must be antisymmetric.

Sometimes we are interested in ordering scientific theories relative to one another. We can do this over time, by successional reductions, or up and down the scientific pyramid, via inter-level reduction. This will be the focus of §5.2. From here on, I will refer to reflexivity, antisymmetry, and transitivity as *ordering relations*, understanding that there are other relational properties that may successfully create an order-like relation amongst members of a set, but that these three are suitable for our purposes.

These three ordering relations are important because I feel that each, pre-theoretically, is a feature that a successful theoretical model of reduction could be interested in preserving. Each of the ordering relations is linguistically pertinent to how we might use “reduction” in scientific parlance. Reflexivity is an interesting case that tests the logical consequences of a reduction archetype against an often unconsidered possibility of having a theory be reducible to itself. Antisymmetry is relevant as most intuitions prevent a reduction from going “either way”, so it is prudent to see if any of our models allow this possibility. Transitivity is perhaps most important, as often questions of inheritance will appear in both intra-level and inter-level reductions. Together the three properties are crucial for telling us about what sort of theories have the potential to reduce to others, intuitively or formally.

§5.1.1 Nagel's Deductive Model:

The first model that we examined in §1.1.1 was the deductive model of Nagel. Recall that there were two types of reductions: homogeneous and heterogeneous. The former concerned theories that shared vocabularies, whereas the latter applied when there were terms that the reduced theory used, yet the reducing theory did not. I previously extracted the following definitions for each case:

β is *homogeneously reduced* to α iff: A derivation of β is possible from α .

β is *heterogeneously reduced* to α if:

- (I) The theoretical vocabulary of β contains terms not in the theoretical vocabulary of α .
- (II) Each term of β is linked to a term or compositions of terms in α by means of well-established biconditionals.
- (III) The derivation of β by means of these biconditionals follows from α .

Before discussing which ordering operations obtain for Nagel's model, I will need to add a bit more detail to the existing model. First, I will add subscripts to each of the conditions so as to indicate the direction of the reduction. For instance, $(I_{\beta-\alpha})$ will represent the condition (I) when attempting to reduce β to α , while $(I_{\alpha-\beta})$ will represent the condition (I) when attempting to reduce α to β . Also, I will restrict the domain of theoretical vocabulary to single-placed predicates. As such, I will specify how (II) should appear only for single-placed predicates, signifying this as (II)^P. I have ignored mention of other types of terms, such as individuals, functions, etc., for I feel that restricting the domains of the theories to predicates will not limit the force of my discussion. In fact, I believe just the opposite is true, for it is not difficult to envision how (II) could be explicated for the other types of terms, and that the proofs and constructions below would simply mimic the reasoning done on single-placed predicates in a slightly different fashion.

Now I will restate Nagel's formal model in more detail:

β is *heterogeneously reduced* to α if:

($I_{\beta-\alpha}$) The theoretical vocabulary of β contains terms not in the theoretical vocabulary of α .

($II_{\beta-\alpha}$)^P For each single-placed predicate B in β that does not occur in α there is a biconditional $(x)(Bx \leftrightarrow F_{\alpha}(x))$ such that:

(i) F_{α} is a composition of one or more predicates of α .

(ii) The biconditional is well-established.

($III_{\beta-\alpha}$) The derivation of β by means of these biconditionals follows from α .

When discussing the ordering relations for Nagel, I will distinguish between the homogeneous and heterogeneous cases. Little space need be wasted when talking of homogeneous reductions, as they very easily satisfy our three relational properties. In homogeneous cases reflexivity occurs, as trivially $\alpha \vdash \alpha$. Antisymmetry is maintained, as any case in which $\alpha \vdash \beta$ and $\beta \vdash \alpha$ will force $\alpha = \beta$, as this is the definition of logical equivalence. Lastly, transitivity is likewise academic, as logical derivations with mutual vocabularies are themselves transitive.

More interesting are the cases of heterogeneity:

(a) Reflexivity

First, reflexivity is not a question that applies to the heterogeneous model: condition ($I_{\alpha-\beta}$) is sufficient to guarantee that reflexivity will never hold; α cannot contain terms that are not in α . The homogeneous version has no restrictions, and we know that it is always possible to derive α from α . Thus Nagel's model is reflexive.

(b) Antisymmetry

With a little elaboration it should be clear that there are a good number of cases where, if β is reduced to theory α , we can still have α reduce to β .

Let's presume that β is reduced to α . As long as we have at least one term from α 's vocabulary that isn't in β , then we are granted ($I_{\alpha-\beta}$). Keeping in mind that we know ($I_{\beta-\alpha}$), this

will happen whenever we find that the terminological language of α is not embedded within the language of β . This is not too restrictive, for otherwise α might be better-termed a *subtheory* of β , and rarely do we seek to establish subtheory-theory reductions⁸⁶.

Now look at the biconditionals stipulated by $(\Pi_{\beta-\alpha})^P$. For each term in the vocabulary of β there must be a biconditional that relates this term to some other term in α , or to some composition of α terms $F_\alpha(x)$. So imagine an α that contains two predicates that a β does not have, A_1, A_2 , and that our theory β contains just four predicates that α lacks, B_1, B_2, B_3, B_4 . We are then granted the following biconditionals:

$$\begin{aligned} (x)(B_1x &\leftrightarrow F^1_\alpha(x)) \\ (x)(B_2x &\leftrightarrow F^2_\alpha(x)) \\ (x)(B_3x &\leftrightarrow F^3_\alpha(x)) \\ (x)(B_4x &\leftrightarrow F^4_\alpha(x)) \end{aligned} \tag{5.1}$$

Let's further stack our deck by explicitly providing each of the four functions:

$$\begin{aligned} F^1_\alpha(x) &: (A_1x \ \& \ A_2x) \\ F^2_\alpha(x) &: (A_1x \ \& \ \sim A_2x) \\ F^3_\alpha(x) &: (\sim A_1x \ \& \ A_2x) \\ F^4_\alpha(x) &: (\sim A_1x \ \& \ \sim A_2x) \end{aligned}$$

Now inserting these components yields the following biconditionals.

$$\begin{aligned} (x)(B_1x &\leftrightarrow (A_1x \ \& \ A_2x)) \\ (x)(B_2x &\leftrightarrow (A_1x \ \& \ \sim A_2x)) \\ (x)(B_3x &\leftrightarrow (\sim A_1x \ \& \ A_2x)) \\ (x)(B_4x &\leftrightarrow (\sim A_1x \ \& \ \sim A_2x)) \end{aligned} \tag{5.2}$$

⁸⁶ Kemeny and Oppenheim see these “internal” reductions as being the concern of J. H. Woodger (1952), and present a definition extracted from his work alongside Nagel’s. I will not deal with derivational “internal” reductions, for Woodger requires that the theoretical vocabulary of α be a proper subset of the theoretical vocabulary of β . The difficulties that lie in this assumption will be showcased in §5.1.2, when Kemeny and Oppenheim themselves presuppose it.

To show that α is reduced to β , we need to be assured that each term in α may be explicated solely in the language of β . But the four relations that are given by [5.2] allow us to write A_1 and A_2 in terms of β -vocabulary as follows:

$$\begin{aligned} (x)(A_1x &\leftrightarrow (B_1x \vee B_2x \vee \sim B_3x \vee \sim B_4x)) \\ (x)(A_2x &\leftrightarrow (B_1x \vee \sim B_2x \vee B_3x \vee \sim B_4x)) \end{aligned} \quad [5.3]$$

Here [5.3] follows from [5.2]⁸⁷. Now if the two biconditionals given in [5.2] are “well-established”, then we are naturally guaranteed that [5.3] will be well-established, as the deduction will certainly preserve endorsement of a scientific community. Notice that [5.3] takes the form:

$$\begin{aligned} (x)(A_1x &\leftrightarrow F^1_{\beta}(x)) \\ (x)(A_2x &\leftrightarrow F^2_{\beta}(x)) \end{aligned} \quad [5.4]$$

Thus we have proven $(II_{\alpha-\beta})^P$.

Lastly, to establish $(III_{\alpha-\beta})$, imagine that the content of each respective theory consists in sentences containing scientific terminology and logical connectives. This is not overly restricting, as it isn't much of a stretch to presume any mathematical knowledge to be common to each vocabulary and theory. Now any theoretical posit of our theory α may be translated into a theoretical posit of our β . This is because a derivation of one theory from another would proceed smoothly via the granted biconditionals along with the derivations provided by $(III_{\beta-\alpha})$.

Thus I have shown that given a reduction of β to α , it is possible to provide conditions that will also allow α to reduce to β . As such a result is significant, I will take some time here to examine its consequences.

We can imagine a story that could possibly be represented by these variables. Let two scientists meet up to discuss their work in a world where the dubious divide between “right-brain people” and “left-brain people” holds true in a curious way. One scientist, a psychologist, has identified three personality types in people: “normal”, “art-incapable” and “math-incapable”. Here the art-incapable people simply cannot appreciate, discuss, or participate in all things

⁸⁷ The explicit derivations are left to the reader.

artistic. Likewise, the math-incapable people are hopelessly unable to reason quantitatively. The last “type of person” that is necessary for the story to work out will be those that are “brain-dead”, rather uninteresting psychological subjects. Now, let the other scientist in the discussion be a neuroscientist. The neuroscientist is attentive to when each of the brain hemispheres are functional, as sometimes one cannot function due to accident, surgery, or perhaps birth defect. Here are each of the terms mapped onto the structure employed earlier:

B₁τ: τ is a normal type

B₂τ: τ is an art-incapable type

B₃τ: τ is a math-incapable type

B₄τ: τ is a brain-dead

A₁τ: τ has a functional right brain hemisphere

A₂τ: τ has a functional left brain hemisphere

During conversation, perhaps the scientists became aware of the following correlations, as I substitute for each of the biconditionals from [5.2]:

(x)(B₁x ↔ (A₁x & A₂x)): A person is normal just in case: they have a functional right brain hemisphere and a functional left brain hemisphere.

(x)(B₂x ↔ (A₁x & ~A₂x)): A person is math-incapable just in case: they have a functional right brain hemisphere and don’t have a functional left brain hemisphere.

(x)(B₃x ↔ (~A₁x & A₂x)): A person is art-incapable just in case: they do not have a functional right brain hemisphere and have a functional left brain hemisphere.

(x)(B₄x ↔ (~A₁x & ~A₂x)): A person is brain-dead just in case: they do not have a functional right brain hemisphere and do not have a functional left brain hemisphere.

Now it is possible all of the theoretical discourse from the psychologist could be recast as theoretical discourse of lacking or possessing functional right/left brain hemispheres. A

psychological observation-regularity such as “art-incapable people tend to get nosebleeds more often than normal or math-incapable people” could be reworked using neuroscientific terminology: “People whose left brain hemispheres do not function typically/normally tend to get nosebleeds more often than those with functional left brain hemispheres”. Granting me a few more biconditionals, we could likely arrive at this more technical substitution: “rupture of the sphenopalatine artery is more likely to occur when a subject’s left brain hemisphere has been inhibited than when it is operational”. Thus for this limited domain of discourse, a reduction of psychology to neuroscience would be achieved.

But recall that we also have enough ammunition to supplant each of the neuroscientific terms for psychological ones, since from [5.2] we may derive [5.3]. Here we have:

$(x)(A_1x \leftrightarrow (B_1x \vee B_2x \vee \sim B_3x \vee \sim B_4x))$: A person has a functional right brain hemisphere if and only if that person is either: normal or art-incapable or not math-incapable or not brain-dead.

$(x)(A_2x \leftrightarrow (B_1x \vee \sim B_2x \vee B_3x \vee \sim B_4x))$: A person has a functional left brain hemisphere if and only if that person is either: normal or math-incapable or not art-incapable or not brain-dead.

The last required stipulation is that each neuroscientific law and conjecture could be shown to be a psychological one, given that the remainder of the terminology is commonly employed in each respective science (and in the same ways). Thus a neurological law stating that “a functional right brain is necessary to remember faces” could be reread as claiming “only normal or math-incapable people can remember faces”. So long as this can be done for each neuroscientific claim that employed only the abovementioned unshared terminology, we would have a reduction of neuroscience to psychology.

It turns out that I have been very careful in the construction of my example: many reductions cannot go both ways for the Nagel model. Given a list of sentences of the sort found in [5.1], a good number of such α -functions chosen to relate each term in β would not permit the derivation of the similar β -functions for every specific term in α . This is because most plausible cases of β to α reduction feature an α fundamentally more terminologically rich than the

corresponding β . This is not just a comment about the number of terms in α , but of their structure when compositionally related to the terms of β . For example, contrast nuclear models that use protons, neutrons, and electrons as basic components with models that instead use quarks of differing generations. Here the quark models are richer in terms of theoretical constituents as well as structural interrelations between these constituents.

(a) Transitivity

Assume that we have a reduction from γ to β and also from β to α . Assume that all three theories are distinct, so that $(\mathbf{I}_{\gamma-\alpha})$ holds. Now take a predicate in γ , $\mathbf{G}$, that is tracked by the grouping $\mathbf{F}_\beta$ to terms from β , $\mathbf{B}_1, \mathbf{B}_2, \dots, \mathbf{B}_n$. In turn, each of these $\mathbf{B}_i$'s has a corresponding grouping $\mathbf{F}_\alpha^i$ that biconditionally links them to some logical composition of terms from α , $\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_m$. We now have all the resources available to construct biconditionals that link each $\mathbf{G}$ to terms that are wholly in α . Take every $\mathbf{B}_i$ that occurs in the biconditional $(\mathbf{x})(\mathbf{Gx} \leftrightarrow \mathbf{F}_\beta(\mathbf{x}))$ and substitute for each the logical composition of $\mathbf{a}_j$'s that is described by the respective composite sentence $\mathbf{F}_\alpha^i$. The end result will be a biconditional that links $\mathbf{G}$ with some function of terms that reside wholly in α . Thus we have satisfied $(\mathbf{II}_{\gamma-\alpha})^P$, so long as we additionally assure that the “well-establishment” clause is transitive.

A theory is established by the consent of competent scientific practitioners: were they to judge α as performing better than β by some criterion, and judge β as performing better than γ by a similar criterion, then the criterion would likewise easily prefer α to γ . To switch criteria when judging theories that are speaking of similar phenomena seems straightforwardly undesirable. This is because I take “judging a theory adequate” to be a tacit expression of the empirical success of a theory. Any biconditional that required establishment by an empirical endeavor will do so by practitioners who make measurements by summarizing data and making note of potential errors. I feel that the method that one employs to scrutinize these data may well change over time. But given that any assessment is more or less contemporaneous, employing similar techniques and experimental tools, it would be bad method to assess one theory by a different empirical standard. Each should be scrutinized equally, else we would be guilty of privileging some theories over others *before* they are tested empirically.

Condition (III _{γ - α}) is much more easily-satisfied. α , [all α - β bridge principles] $\vdash \beta$. Similarly, we also know that β , [all β - γ bridge principles] $\vdash \gamma$. Thus we can see that, together: α , [all α - β bridge principles], [all β - γ bridge principles] $\vdash \gamma$. And recall that the bridge principles stack transitively, so we now substitute to find that: α , [all α - γ bridge principles] $\vdash \gamma$. This shows that the relation is transitive, granting that all three theories are distinct.

Does not hold for *every* such α , β , and γ ? Consider a case where we seek to show that: if β reduces α and α reduces β , then α reduces α . Here the fact that some instances of asymmetry are allowed, and the fact that reflexivity is never allowed for heterogeneous cases, readily leads to the conclusion that we cannot in all cases have transitivity. It is correct to use the heterogeneous model to treat a α - β reduction, but not an α - α reduction. The latter should use the homogeneous model, where reflexivity always occurs. Thus we can still consider the relation transitive in every case.

§5.1.2 Kemeny and Oppenheim's Disjoint-Explanation

Model:

As discussed in §1.1.2, Kemeny and Oppenheim provide a model that focused on how well theories explain the available data. Unlike Nagel, Kemeny and Oppenheim formally present their reductive criteria:

Red(α , β , O) [a reduction of β by α relative to observational data O], if:

- (I) The theoretical vocabulary of β contains terms not in the theoretical vocabulary of α .
- (II) Any part of O explainable by means of β is explainable by α .
- (III) α is at least as well-systematized as β . (Kemeny and Oppenheim 1956, 13)

We will rely on the above definition, and also the notion of an “internal reduction”.

Understanding **Voc(δ)** to mean *the vocabulary of theory δ* , the following definition is helpful:

Definition 5: Intred(α , β) [an internal reduction of β by α], if and only if:

Red(α , β) and Voc(α) is a proper subset of Voc(β)
(Kemeny and Oppenheim 1956, 14)

As mentioned previously in §1.2, the concept of an internal reduction does not cleanly translate into either of the categories demarked by inter-level or intra-level reductions. As it will not take too much additional space, what follows will consider our ordering relations both for reductions and for internal reductions⁸⁸.

(a) Reflexivity

(I) requires some terms to differ, clearly disallowing reflexivity. However, when we consider the less-referenced case of an internal reduction, we can come close. If *Definition 5* required simply a *subset*, rather than a *proper subset*, this would allow for reflexivity in all cases. So in either case we cannot have a reflexive reduction.

(b) Antisymmetry

Note that (I) is the same as in Nagel's heterogeneous model, because Kemeny and Oppenheim also require that some terms be exclusive to β . More troublesome is (III), as it should be clear that "systematization" is in need of a more precise definition if it is a quality that is to be compared. Lawrence Sklar notes that for Kemeny and Oppenheim "the asymmetry is introduced through the notion of the systematic power of a theory" (Sklar 1967, 114). As we noticed earlier, Kemeny and Oppenheim have intentionally left the terms "systematic power" and "at least as well systematized" vague, but any attempt to make the concept more precise, I think, is bound to fail. Let us say that systematizing a theory is the introduction of structure to a theory – to make it into a well-defined web of connected parts. Further let the systematization of theory present a well-organized summarizing method for representing all

⁸⁸ Oppenheim and Putnam define a "micro-reduction" (1958, 7). Without going into too much detail, this is an inter-level reduction definition that seeks to account for compositionality. In their discussion, they conclude that micro-reductions are always transitive, never reflexive, and always asymmetric. Our results correlate, with the exception of the asymmetry. This is understandable, due to the prohibition that "the objects in the universe of discourse of β are wholes which possess a decomposition into proper parts all of which belong to the universe of discourse of α " (Oppenheim and Putnam 1958, 6).

possible observation statements such a theory would seek to endorse. From this vantage point, it is quite possible that we would find phylogenetics to be “better systematized” than many other theories. The manner in which modern biology classifies and categorizes the massive spread of evolutionary progress might be seen as superior to similar treatments of other scientific areas. For example, when classifying different states of matter, there has been difficulty categorizing intermediary phases (Jaeger 1998). This would then seemingly prevent any reductions of evolutionary biology from the standpoint of thermodynamics or any other theory that relied on a problematic schema for classifying phases. But perhaps in some cases we could claim that chemical theory, with transition rules and the periodic table, does this job *just as well*. Standard model physics, likewise, could be understood to be a bit less systematized; however if we wave our hands a bit and admit that the nice clean tables of elementary particles, their group structure, and a host of equations do about as good a job being systematized as any of the other discussed theories, this would allow the reduction to happen. But notice that if both theories are equally systematized, this is precisely what would be required to make a reduction symmetric. So in allowing a loose equality in degrees of systematization, symmetric reductions are not excluded.

Kemeny and Oppenheim worry that the condition placed by (III) might thus be too weak. Their suggestion is to amend (III) so as to change “is at least as well systematized” to “is better systematized” (1956, 12). The problem remains that by this “strengthening”, they undermine any possibility of accomplishing the biology-to-chemistry-to-physics chain, for reasons discussed above. If we are to view (III) in this light it will surely prevent any reductions from being symmetric. But to strengthen (III) in this manner will come at the price of preventing most desired antisymmetric reductions from happening as well. So I think that little is to be accomplished by reworking (III) – indeed this is because I do not think that Kemeny and Oppenheim can produce a satisfactory account of “systematization” that will accomplish what they had intended.

So, in granting that there is little way to fix (III), I think that we are led to conclude that, as with Nagel, the reductive account presented by Kemeny and Oppenheim also admits some symmetric reductions. Notice how none of these cases can be internal reductions, as the only way we could have both $\text{Voc}(\alpha)$ be a subset of $\text{Voc}(\beta)$ and $\text{Voc}(\beta)$ be a subset of $\text{Voc}(\alpha)$ would be if they were not proper subsets, a condition which the authors expressly disallow.

(a) Transitivity

Presume that α reduces β and β reduces γ . Now assume that $\alpha \neq \gamma$, as when $\alpha = \gamma$ we would be assuming symmetry and then checking to see if it was reflexive – a case which we have already determined cannot occur. So by letting α and γ differ we have two options: the two theories have different vocabularies (thus assuring (I)), or they are two theories that employ the same vocabularies in different manners. The latter case prevents there from being a reduction, as we observed in our discussion of transitivity above in §1.1.1. Additionally, concerning internal reductions, we would also be allowed to assume that $\text{Voc}(\alpha) \subset \text{Voc}(\beta)$, and that $\text{Voc}(\beta) \subset \text{Voc}(\gamma)$. As transitivity holds for proper subsets, we would know that $\text{Voc}(\alpha) \subset \text{Voc}(\gamma)$.

Now take any portion of an observational set $\mathbf{O}$ that γ can explain. We know that β can explain this piece just as well, as we knew such a fact from the assumptions. And then if β can explain it, we are granted that α likewise explains it as we are to assume that α reduces β . Thus (II) is easily satisfied, as α explains all of any portion of $\mathbf{O}$ that γ does. Finally we have an issue of systematization when looking for (III). It turns out, even without being generous with comparisons involving the term “systematize” (as we were when considering asymmetry), that a more restrictive interpretation is still readily transitive. To be more systematized is to have the capacity for a more discerning interface with the data, as well as having superior organizational and structural features; in either case, it should be apparent that if β is as well systematized as γ , and that α is as well systematized as β , that α will be systematized to at least the same degree as γ . As each case holds, we find that Kemeny and Oppenheim’s reductive account maintains transitivity. Furthermore this same discussion holds for internal reductions, allowing us to conclude internal reductions are transitive as well.

§5.1.3 Suppes’s Semantic-Isomorphism Model:

Suppes’s conception of reduction, discussed in §1.1.3, required the following rule:

(I) For every $\mathbf{M} \models \beta$, there exists an $\mathbf{M}' \models \alpha$ such that $\mathbf{M} \approx \mathbf{M}'$.

Here $\approx$ represents an isomorphism, sample rules for which can be found in §1.1.3. Schaffner proved that for every case in which an α , β satisfy the requirements of a Nagelian reduction (as extracted and formalized by Schaffner), these α , β will also satisfy (I) above (1967, 138). This result will not influence our result for two reasons: first, since my extraction of Nagel's conditions slightly differs from Schaffner's, I would need to reprove the theorem⁸⁹. Second, Schaffner's result is a conditional, and for us to be able to import results from §5.1.1 it would need to be a biconditional.

Lastly, before we proceed, it is important to remember that (I) was likely intended only as a necessary condition for a reduction. What follows should therefore be regarded as discerning necessary conditions for our three target ordering properties.

(a) Reflexivity

Suppes' reduction model allows for reflexivity straightforwardly: a model of a theory α need not be extended, as it is trivially automorphic.

(b) Antisymmetry

Let us grant that we have α reducing β . We now know that every $\mathbf{B} \models \beta$ has an extension $\mathbf{B}'$ that will be isomorphic to some $\mathbf{A} \models \alpha$. Now suppose that $\alpha \neq \beta$, and that there is an instance when a $\mathbf{B}$ can only be fit to a model $\mathbf{A}_1$ that is larger than $\mathbf{B}$. Put another way, require that in one case it must be that $|\mathbf{A}_1| = |\mathbf{B}'| > |\mathbf{B}|$. Examine this case when trying to show that β reduces α . Here we must be able to show that $\mathbf{A}_1$ can be mapped on to some $\mathbf{B}$. But, as we are only allowed to extend the models of the reduced theory, we know that there is no $\mathbf{A}_1'$ that would suffice to be isomorphic to any of the $\mathbf{B}$'s. This is because: $|\mathbf{A}_1'| \geq |\mathbf{A}_1| > |\mathbf{B}|$. Thus, it must be that in all cases $|\mathbf{A}'| = |\mathbf{B}|$.

The models of the two theories must be equicardinal if there is to be symmetry at all. Furthermore, each of the predicates, functions, and name letters must match to ones with the exact same assignments. They will at best be labeled differently, but structurally they are

⁸⁹ As hinted at in §1.1.3, this could be done without much trouble.

identical. In this equicardinal, identical assignment case, we will have α reduce to β , without extension. But if one *only* looks at the structure that each theory forces on its objects, as the structuralist is wont to do, we find that α and β are *by definition* identical. So by the structuralist's focus on the semantic and structural entailments of the theories, we can show that the two theories whose models are equicardinal and isomorphic are *the same theory*. The only cases where we are allowed to have both α reducing β and β reducing α is the case where $\alpha = \beta$, which is precisely our definition for antisymmetry.

(c) Transitivity

Assume that β reduces γ . For any model $G \models \gamma$, we know that we may always extend G to create a new model G' that will be isomorphic to a model $B \models \beta$. We also know that α reduces β , and thus that for every $B \models \beta$ there exists an extension B' such that for some $A \models \alpha$, $A \approx B'$. Take a model G that satisfies γ . Then when we extend the model to G' , recall that we merely add components while the kernel of G remains. We know that such a model is isomorphic to B , and that we may likewise extend this to B' . In this fashion we are reminded that every generated B' contains a kernel that is isomorphic to G . So we may merely extend G directly into a G'' , where the $G'' \approx B'$, and thus we are assured that $G'' \approx A$. Hence the Suppes reductive relation is transitive.

§5.1.4 Schaffner's Model:

In §2.1, we discussed how Schaffner altered the model of Nagel to accommodate the difference between theories. For our purposes, we will rely on the following formulation, a combination of (Schaffner 1967) and (Schaffner 1974):

β is *reduced* to α iff:

- (I $_{\beta-\alpha}$) The individuals/predicates of β^* are linked with individuals/predicates or groups of individuals/predicates of α by empirically-supported reduction functions.
- (II $_{\beta-\alpha}$) β^* can be derived from α by means of these reduction functions.

(III _{$\beta-\alpha$}) β^* corrects β , in the sense that:

(i) β^* provides more accurate predictions than β in most cases.

(ii) β^* explains why β was incorrect, and why it was once considered correct.

(IV _{$\beta-\alpha$}) β^* and β are linked by a relation of strong analogy, [A_s].

(a) Reflexivity

The goal of Schaffner's model was to accommodate reductions between theories that differed in predictions, with the reducing theory doing a better job than the reduced at this task. The reductive criterion reflects this in each of the portions of the third clause. When seeking to provide an α - α reduction, we would not need to add any reduction functions, because all vocabulary is mutual. This makes (I _{$\alpha-\alpha$}) and (II _{$\alpha-\alpha$}) trivially satisfied, as we let α^* be the same as α . Now (III _{$\alpha-\alpha$})(i) would read " α^* provides more accurate predictions than α in most cases", which would become " α provides more accurate predictions than α in most cases" – a condition that cannot happen. Similar confusion for (III _{$\alpha-\alpha$})(ii), which would ultimately read: " α explains why α was incorrect, and why it was once considered correct". It is difficult to see how a genuine scientific theory would have the potential to explain its own incorrectness, and furthermore do so correctly. As a result, reflexivity cannot occur.

(b) Antisymmetry

The majority of possible symmetric cases will have a similar worry. First by (III _{$\beta-\alpha$})(i) it must be that " β^* provides more accurate predictions than β in most cases", and from (III _{$\alpha-\beta$})(i) it is required that " α^* provides more accurate predictions than α in most cases". In some ideal cases, the analogy merely changes vocabulary and trivial correspondences. Here the predictions of an analog theory are exactly those of the reducing theory, yet they are now couched in the terminology of the reduced. Assuming that the reduction functions were adequate, we now realize that at best the predictions of α^*/β^* will differ from those of β/α only terminologically, not in what they require of the empirical. Assume that we have such an analogy, allowing us to employ both α^*/β^* and α/β interchangeably for purposes of determining the accuracy of predictions. This will generate the following two incompatible propositions: " α provides more

accurate predictions than β in most cases”, and “ β provides more accurate predictions than α in most cases”. Thus as long as $\beta^*[A_s]\beta$ and $\alpha^*[A_s]\alpha$ are strongly analogous in the above sense, there can be no symmetric cases.

Could we then create a β^* and α^* so that each made better predictions? I submit that we could, but that this would be a case where the analogy between them, $[A_s]$, ceased to be a strong one. If each analog theory deviated enough so as to swing the majority of empirical successes in favor of the other, then I contest that one of the analog theories would differ too much from its original. The requirements of $(III_{\beta-\alpha})(i)$, when taken in consideration with the dictate that the $[A_s]$ be strong, disallow any symmetric cases.

(c) Transitivity

Transitivity occurs in some, but not all, cases. I will provide a case where it does hold, and then modify this case to provide an example where it does not hold. For our case let each theory, α , β , γ , contain only one equation as follows:

$$\alpha: a(x) = x$$

$$\beta: b(x) = x+1$$

$$\gamma: g(x) = x+2$$

Let the three theories differ in their vocabulary, allowing a straightforward substitution by reduction functions. To do so in predicate logic, let it be that:

$$A\tau: \tau \text{ is an } a.$$

$$B\tau: \tau \text{ is a } b.$$

$$G\tau: \tau \text{ is a } g.$$

Now relate these by reduction functions, which comprise $(I_{\beta-\alpha})$ and $(I_{\gamma-\beta})$:

$$F_{\beta-\alpha}: (x)(Ax \leftrightarrow Bx)$$

$$F_{\gamma-\beta}: (x)(Bx \leftrightarrow Gx)$$

When combined, the theory and the translation function that correspond to it will allow for the derivation of analog theories:

$$\begin{aligned} (\alpha \ \& \ F_{\beta-\alpha}) &\rightarrow \beta^*, \text{ where } \beta^*: \mathbf{b}^*(\mathbf{x}) = \mathbf{x} \\ (\beta \ \& \ F_{\gamma-\beta}) &\rightarrow \gamma^*, \text{ where } \gamma^*: \mathbf{g}^*(\mathbf{x}) = \mathbf{x}+1 \end{aligned}$$

Each of these entailments are given as $(\mathbf{II}_{\beta-\alpha})$ and $(\mathbf{II}_{\gamma-\beta})$.

This discussion allows us to create a reductive function between α and γ :

$$\mathbf{F}_{\gamma-\alpha}: (\mathbf{x})(\mathbf{Ax} \leftrightarrow \mathbf{Gx})$$

$\mathbf{F}_{\gamma-\alpha}$ is a consequence of $\mathbf{F}_{\beta-\alpha}$ and $\mathbf{F}_{\gamma-\beta}$. Thus, as $\mathbf{F}_{\beta-\alpha}$ and $\mathbf{F}_{\gamma-\beta}$ are empirically corroborated, so will be $\mathbf{F}_{\gamma-\alpha}$. It could well be that through different levels of organization, it was uncontroversial that $\mathbf{a}$'s were $\mathbf{b}$'s, and $\mathbf{b}$'s were $\mathbf{g}$'s. Thus the jump to claiming $\mathbf{a}$'s were $\mathbf{g}$'s follows with equal empirical weight. In this way we can claim that $(\mathbf{I}_{\gamma-\alpha})$ is satisfied.

To establish $(\mathbf{II}_{\gamma-\alpha})$, we must first remember that α in combination with $\mathbf{F}_{\gamma-\alpha}$ will lead to an analog theory for γ , but that this need not be identical to the above γ^* . Instead we will refer to it as $\gamma^{*'}$, as it will differ:

$$(\alpha \ \& \ F_{\gamma-\alpha}) \rightarrow \gamma^{*'}, \text{ where } \gamma^{*'}: \mathbf{g}^{*'}(\mathbf{x}) = \mathbf{x}$$

By assumption $(\mathbf{III}_{\beta-\alpha})(\mathbf{i})$ and $(\mathbf{III}_{\gamma-\beta})(\mathbf{i})$ happen; let this be the case on grounds that the actual measured values of the phenomena in general are very close to $\mathbf{x}$. Therefore, for all values $\mathbf{g}(\mathbf{x}) > \mathbf{g}^*(\mathbf{x})$ and $\mathbf{b}(\mathbf{x}) > \mathbf{b}^*(\mathbf{x})$. Thus we are assured that $\mathbf{g}^{*'}(\mathbf{x})$ will be a better measure than $\mathbf{g}$, as we will always have $\mathbf{g}(\mathbf{x}) > \mathbf{g}^{*'}(\mathbf{x})$, showing $(\mathbf{III}_{\gamma-\alpha})(\mathbf{i})$. $(\mathbf{III}_{\beta-\alpha})(\mathbf{ii})$ and $(\mathbf{III}_{\gamma-\beta})(\mathbf{ii})$ will also come with our assumptions as explanations of the inadequacy of the β by β^* and γ by γ^* . To provide an explanation of the inadequacy of $\gamma^{*'}$, simply mention all that was involved in the explanations given by β^* of β and γ^* of γ . This will provide enough for $\gamma^{*'}$ to explain the past successes and current lack of success of γ , for $\gamma^{*'}$ is isomorphic to β^* . So we also are guaranteed of $(\mathbf{III}_{\gamma-\alpha})(\mathbf{ii})$.

Last, it remains to show that the difference between the analog theory and the reduced theory is satisfactory – we must prove $(\mathbf{IV}_{\gamma-\alpha})$. Because I have provided a simple case where our

theories each consist in one single equation, any analysis of “closeness” will be merely quantitative. Borrowing some terminology from §3.4.1, we might do this in terms of a validity limit, δ . Imagine that $[A_s]$ is defined in such a way that a difference of more than 3 was unacceptable, that is, when $\delta < 3$. As required, the differences between predictions of β^* and β as well as γ^* and γ are acceptable, as $\delta=1$ for each case. Between $\gamma^{*'} and γ , $\delta=2$, so we find that $[A_s]$ is also acceptable here. $(IV_{\gamma-\alpha})$, the last piece needed to show a transitive case, occurs. The Schaffner model allows some cases where α reducing β , and β reducing γ , further assures that α reduces γ .$

Recall that the Schaffner model was interested in letting reductions occur where one theory approximated another. Thus to show how transitivity may fail, I will construct a case where the jumps from α to β , and β to γ are acceptable, yet the jump from α to γ is not. To show a case that is intransitive, leave everything the same in the past example except the requirements on $[A_s]$. Make it that an acceptable difference is now $\delta < 1.5$. Now this will leave $(VI_{\beta-\alpha})$ and $(VI_{\gamma-\beta})$ untouched, but disallow that $(IV_{\gamma-\alpha})$ happens. This is because a $\delta=2$ difference will no longer meet our more stringent $[A_s]$ requirements. So there is a case where the Schaffner model allows that α reduces β , and β reduces γ , yet α does not reduce γ .

§5.1.5 The New Wave Model:

In the Schaffner model, the analog theory was a consequence of α , yet it was formulated from the theoretical resources of β . That is why we had referred to it as β^* . The New Wave model detailed in §2.1.1 instead formulates the analog from α by constraining it under conditions C , which I will here represent⁹⁰ as α^C . The difference here is that the analog theory is now constructed from within α 's framework, making α^* a better descriptor. Indicating the analogy relation by $[A_n]$, I repeat the New Wave's reduction criterion, restating [2.8]:

$$\beta \text{ reduces to } \alpha \text{ iff: } \alpha^C \rightarrow \alpha^*, \text{ and } \alpha^*[A_n]\beta \quad [5.5]$$

⁹⁰ Another possible formulation is to conjoin C with α , represented as $(\alpha \& C)$ in [2.7]. As I feel [5.5] (identical to [2.8]) is more general, I will employ this interpretation instead.

(a) Reflexivity

Let C be null, causing $\alpha^C = \alpha$. Also let $\alpha^* = \alpha$. Now we can see that under these constraints every theory will be self-reducible, for trivially $\alpha^C \rightarrow \alpha^*$ as $\alpha \rightarrow \alpha$. There is no worry about the satisfaction of $[A_n]$, as this case demands that $\alpha[A_n]\alpha$ – a maximally strong analogy because it is a relationship of identity. None of this should come as much of a surprise, for it is quite easy to make “an analogous theory of itself from itself”.

(b) Antisymmetry

There are cases where symmetry does not hold: the case study from §2.2 cannot proceed both ways. This is due to the fact that the Trautman Recovery Theorem does not proceed cleanly in both directions. However, to provide a case of symmetry, examine the following case involving two theories, α , β , that consist in the following equations:

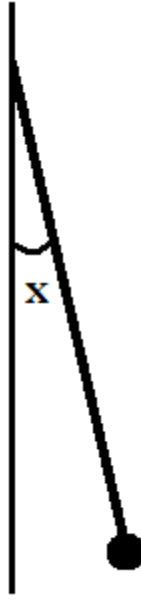
$$\alpha: a(x) = \mu \sin(x)$$

$$\beta: b(x) = \mu x$$

Let the constraint C be “to let the x ’s be quite small”, or using parlance from §3.4.1, “have $x \ll 1$ ”. One further way of expressing C would be: “examine the asymptotic region around the $x \rightarrow 0$ limit”⁹¹. Now we find that α^C leads to an α^* that is indeed very much like β . This is because as we let x approach 0 , $a(x)$ and $b(x)$ approach each other. They appear (quantitatively) analogous, as the observed behavior of each function is quite similar.

Remember that in §3 (most notably §3.3), one lesson was that claims of two functions “approaching one another” are easily satisfied mathematically, and potentially problematic when requiring a physical interpretation. So I will provide empirical circumstances to evaluate our claims. Let $a(x)$ and $b(x)$ be measures of the angular acceleration experienced by an ideal pendulum, as shown by {Figure 5.1}:

⁹¹ This is quite different than the locution “taking the limit as x approaches 0 ”. As we saw in §4, there is a difference between behavior *at the limit* and behavior *around the limit*.



{Figure 5.1}

Here x is the angle the pendulum makes to the vertical rest position, while μ is a constant comprised of initial conditions⁹². For an example of a “small x ”, when $x = 2^\circ$, and $\mu = 10 \text{ rad/s}^2$ corresponding to household initial conditions⁹³, the difference between $a(x)$ and $b(x)$ will be less than .03%. Because this is quite difficult to detect experimentally, I presume that $\alpha^*[A_n]\beta$. Thus we find that β reduces to α .

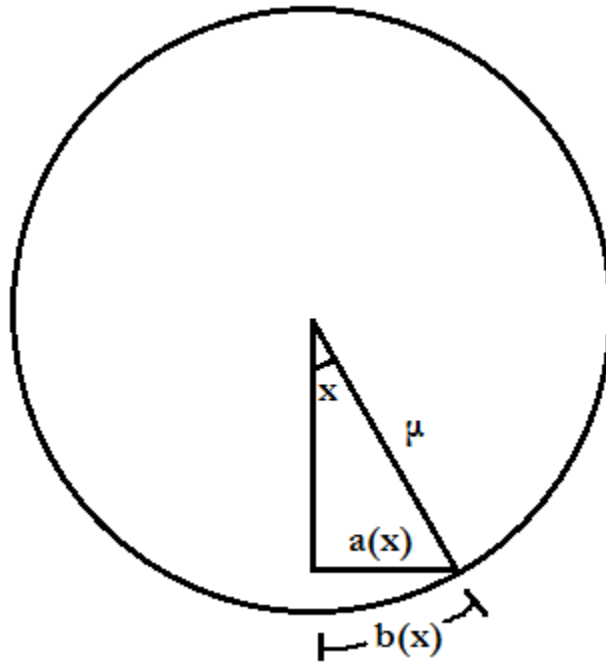
But for the pendulum case we could just have easily begun with β . We would perform the same C to β , requiring that we pay attention only to small pendulum angles. β^* is very much like α under this constraint; in the same way that $\alpha^*[A_n]\beta$ was acceptable, we must also accept $\beta^*[A_n]\alpha$. In doing so we will find that α reduces to β , showing that the above example is a case of a symmetric New Wave reduction.

One objection I can imagine being raised is that, despite the ability to start from β and arrive at α , to do so is simply bad science. The correct measure represented by $a(x)$ has much more “structure”, whereas the simple linear model provided by $b(x)$ is obviously the incorrect one. As we saw for the GR-NM reduction of §2.2, the more complex theory GR served as our α , and only after we had “screened out” some of the complexity by C did we arrive at something

⁹² Specifically, have $\mu = -g/l$, where g is the acceleration due to gravity and l is the length of the pendulum arm.

⁹³ Using $g = 10 \text{ m/s}^2$ and $l = 1 \text{ m}$.

that resembled the more course β of NM. Similarly this was the case for the SR-CM reduction of §3: the linear measure of CM took the role of β .



{Figure 5.2}

To answer the objection, I will provide an alternate scientific case that uses the same equations for α and β , yet where β is the correct measure. Examine the small slice of the circle as shown in {Figure 5.2}. The correct measure for the length of the arc would be to use $b(x)$, where x is the angle and μ is the radius of the circle. Alternately, we could inscribe a right triangle and use the small side's value as the arc length – $a(x)$ under the same interpretation. To create a reasonable scenario, have the circle slices be pizza sized ($\mu = r = .5$ m), and let the angle $x = 10^\circ$. In this case there will be a difference of less than .5% between α and β , quite an acceptable difference for our tabletop scenario. Thus $[A_n]$ will hold in this case, and by the same reasoning as above we may claim that β reduces to α , as well as that α reduces to β . This second example shows that every case need not privilege the theory containing “more structure”.

(c) Transitivity

For there to be a transitive New Wave reduction, we first need two reductions, an β - α reduction and a γ - β reduction:

$$\begin{aligned}\alpha^{C1} &\rightarrow \alpha^*, \text{ and } \alpha^*[A_n]\beta \\ \beta^{C2} &\rightarrow \beta^*, \text{ and } \beta^*[A_n]\gamma\end{aligned}$$

This requires the existence of two analog theories: α^* which bridges α to β , and β^* which bridges β to γ . Once these are in place, we would need to modify α under some conditions **C3** to turn it into an analog theory α^{*} . This analog theory is meant to mimic γ from the conceptual resources of α .

Transitivity can happen in some cases. To provide our example, let each of our three theories focus on the quantitative measure of a variable x as follows, letting k , m , n be constants:

$$\begin{aligned}\alpha: a(x) &= kx + mx^2 + nx^3 \\ \beta: b(x) &= kx + mx^2 \\ \gamma: g(x) &= kx\end{aligned}$$

The conditions, **C1**, **C2**, and **C3**, will in each case be limits. For α we will impose the condition **C1** of “limiting n to 0”, causing α^* to be identical with β . Similarly for β we will impose the condition **C2** of “limiting m to 0”, creating a β^* that is the same as γ . Here there can be little controversy of the “closeness” of the analog theories, so discussion of $[A_n]$ isn’t needed. Also, as an assumption, allow that the limits performed on each of the constants are legitimate and are justified on good physical grounds.

Now it remains to show that it is possible to impose some conditions, **C3**, on α to make it into an analog theory α^{*} , that will stand in relation $[A_n]$ with γ . Here the conditions **C1** and **C2** will provide our guide: let **C3** be “limiting m to 0, and limiting n to 0”. As m and n are mere constants, and being assured that the limits are empirically acceptable, α^{*} will be analogous to γ . Because in the limits $a(x) = g(x)$, the $[A_n]$ is as strong as one could want. Thus our example provides an instance of a transitive reduction, albeit a very simple one.

What would be required for transitivity to fail? One way would be to have the analogous bridges across the β - α and γ - β divides be acceptable, but the γ - α bridge not satisfactory for an

[A_n]. This is quite similar to the scenario that I employed to show an intransitive case for the Schaffner model. For this case we will employ the following equations within each theory:

$$\alpha: a(x) = kx + mx^2 + nx^3 + 2$$

$$\beta: b(x) = kx + mx^2 + 1$$

$$\gamma: g(x) = kx$$

Impose the same constraints as last time:

C1: “limiting n to 0 ”

C2: “limiting m to 0 ”

C3: “limiting m to 0 , and limiting n to 0 ”

Lastly, require that for an [A_n] to be acceptable, a given theory must make predictions that are within $\delta \leq 1$. Now we are assured that our assumptions hold: $\alpha^{C1} \rightarrow \alpha^*$, and $\alpha^*[A_n]\beta$; $\beta^{C2} \rightarrow \beta^*$, and $\beta^*[A_n]\gamma$. This is because the difference between both α^*/β and β^*/γ are within 1 , an acceptable δ for [A_n]. However once we impose **C3** on α to arrive at the analog theory α^{*} , the α^{*}/β difference is too great. Here $\delta = 2$, and as such $\alpha^{*}[A_n]\gamma$ does not hold.

§5.2 Orderings of Science:

Here I consider the influence that the aims of those seeking a reduction should have on the development of philosophical projects that attempt to codify a model for reduction. As stepping stones for this project, I have proven which of the relational features of reduction obtain for the above five archetypes. {Figure 5.3} is a summary of the results of §5.1.

	Reflexivity	Antisymmetry	Transitivity
Nagel	Always	Sometimes	Always
K&O	Never	Sometimes	Always*
Suppes	Always	Always	Always
Schaffner	Never	Always	Sometimes
New Wave	Always	Sometimes	Sometimes

*Excepting cases where $\alpha = \gamma$

{Figure 5.3}

I will now discuss how each relational property could be considered with respect to certain goals one might have for a reduction. My aim is to show what is at stake for reduction with regards to the ordering relations. By doing so we can begin to understand possible motivations that could underlie models of reduction that include any of the above relational properties.

§5.2.1 Reflexivity and Triviality:

If one's reductive goal is ontological, then being able to show reflexivity is vacuous. Were we to find that a theory reduced to itself, it would show that the constituents of the reducing theory were the same as the reduced theory – accomplishing nothing. Similarly, if it were the case that a model disallowed a same-theory reduction, no ontological implications could follow necessarily, for it would not be meaningful to say that “a thing was not composed of the things which it was composed of”, or some similarly confusing locution. Thus those interested in proving ontological claims by means of reductions would be wholly disinterested as to whether a theory could or could not be said to be reducible to itself.

Epistemic goals would seem also to have a hard time making sense of reflexivity: what would it mean for a theory to be able to explain itself? There are worries about odd cases of “self-explaining explainers”, but they would likely not here be applicable even if one were to accept that they genuinely occurred. Instead, I think that it certainly could be the case that one part of a theory was explaining another part of that same theory. One might allow for the relata in a “theoretical reduction” to be components of a theory, as opposed to the entire body of the theory itself. Assume that someone was interested in showing that we might explain equation E_1 of a theory T by another equation E_2 of T through a theoretical reduction. Here it might be tempting to claim that we have gathered an explanation by means of “reducing T to T ”, thereby salvaging a meaning for reflexivity in an epistemically-minded reduction. As attractive as this might sound, I would claim that we have instead produced a particular type of intra-theory reduction, one which is best characterized as “reducing E_1 to E_2 ”. Here the basic elements of the reduction are not “theories” that have many individual constituents, but “theories” understood as representing theory-parts. This answer seems further satisfying when one realizes that reflexivity truly wants “the same thing” reduced to itself, not “different parts of a thing” to be reduced.

Other epistemically-minded goals exist, but in each case I do not think that reflexivity may aid in accomplishing any of them. A theory can do little to explain its own errors by means of a reduction and it cannot transfer confidence to itself. Likewise a theory cannot solidify its own role in the progress of science, for there are no other theories it compares itself to – in a reflexive reduction there is only one relatum. Nor can any of the new goals identified in earlier chapters be realized by a reflexive reduction. This even holds true for the goal of fecundity highlighted in §4.5, as any improvements on a theory in relation to itself are better considered “theoretic improvements” than active attempts at “reflexive reduction”.

In relation to each of these possible goals, I think that we may confidently claim that reflexivity is an aberrant case for reductions. At the onset, it might have appeared worthwhile to either seek to preserve or not preserve reflexivity. This is because for many relations, the reflexive case gives insight into the relation or its relata at some abstract level, or at the very least it provides perhaps a base-case for reductive models and their applicability. Both Kemeny and Oppenheim and Schaffner take the necessary measures to prohibit reflexivity by their definition. It is unclear if this was intentional. However it would have been quite easy for Kemeny and Oppenheim to require it to be a “subset” rather than a “proper subset”. Likewise, Schaffner could

have easily amended either clause of (III_{β-a}) to allow for reflexivity. On the other hand, the addition of another clause or stipulation to a model would have been easy for Nagel, Suppes, or the New Wave to likewise *prevent* the possibility of reflexivity. Since this wasn't done, either each thought it not important enough to warrant consideration, or alternatively, saw that it was important for the self-consistency of reduction in their models. I am tempted to believe it more the former, as there is no discussion of reducing a theory to itself in the literature provided by the authors.

Other goals that we have examined likewise fail to apply. For example the goal of “transferring confidence”, highlighted in §3.3.2, certainly cannot be relevant in a reflexive reduction: the confidence we have in a theory will be the same whether the theory is reducible to itself or not. This is because there is no “transference”, since only one theory is involved. Any goal that pertains directly to successional reductions is likely to fail, for there is no time difference and indeed no “succeeding”. Likewise, we could not classify any reflexive reduction as an inter-level reduction, for there is can be no difference in “level” if the reduction only involves one theory. Any goals that are specific to inter-level reductions would thereby not apply.

I conclude that reflexivity is an unimportant relation for reduction, and only Nibergall disagrees explicitly. When listing which pre-theoretical intuitions we should have about the concept of “reduction”, he mentions that it is important that “each entity is reducible to itself”(Niebergall 2002, 148). I submit that this is a desire imported from highly mathematical conceptions of reduction, ones that treats “reduction” much like the concept of “equivalence”. Reflexivity does little, it seems, either to show self-consistency of a theory or give us confidence in the robustness of a given reductive model. Inasmuch as “reduction” is a concept we seek to apply to science and the activities of scientists, we see no evidence of any scientist interested in finding a “self-reducible theory”. Imagine that we attempted to provide a reduction of GR to GR in §2.2, or of wave optics to wave optics in §4.1 – I find it difficult to see what benefit these would provide. At best reflexivity is a mathematical vestige that does little work in refining the concept of “reduction”, as is vindicated by a consideration of the legitimate philosophical goals one might have for the relation.

§5.2.2 Symmetry and Equivalence:

If one had ontological goals, the possibility of symmetry in a model of reduction would be puzzling. Presume that two theories, **A** and **B**, each make reference to distinct particulars, such as atoms and salt molecules, respectively. If a model of reduction is supposed to underwrite a substantial ontological claim, then if we were able to reduce **B** to **A**, it would give some evidence for claiming that salt molecules were composed of atoms, that salt molecules were nothing but atoms, etc. The problem arises when we are provided with a reduction of **A** by **B**. Here we would be given support for a claim of an “atoms-to-salt-molecules” ontology, something seemingly undermining the work done by the **B** to **A** reduction. Thus antisymmetry seems preferable, for it provides a clear ordering of which entities should have ontological primacy. Even if we are to presume that the ontological insights are mereological, in that they told us about the compositional features of components without any claims to the “ultimate substance” in a comparison, such claims still seem to disallow symmetry.

Explanations that occur for successional reductions typically explain old theories by their successors. Such explanations make lucid that which was otherwise confused or unsupported by the prior theory. Here there is a notion of scientific progress, ideological and historical, that is preserved by the ordering of the reductions. Symmetry would confuse the important temporal arrows that a successional reduction was meant to reinforce. If we were able to have a prior theory explain portions of the successor theory, these explanatory successes might give reason to think that the succeeding theory was inadequate, that it was lacking, or that it perhaps should not be seen as a replacement to the prior theory in the first place.

Inter-level reductions have explanations that are a bit more complicated. What reasons might there be for thinking that α could explain β , but that subsequently β should not be able to explain α , aside from fundamentalist biases? It seems that, even for the fundamentalist, there aren't *a priori* reasons for thinking biology could not at times explain certain mechanisms of chemistry. Before we go down this road, I think that again we may illuminate the problem by paying close attention to precisely *what* is being reduced to *what*. Even if we are to allow that parts of biology may explain parts of chemistry, we should not thereby generalize that the constituents of the reduction are thus in these cases “biology” and “chemistry”. Instead we

should consider how it might be that portions of theories are doing the work as the relata for the α 's and β 's. Assume that we find that a theory-part $\mathbf{a}_1$ of a theory $\mathbf{A}$ reduces an equation $\mathbf{E}_1$ of theory $\mathbf{B}$, thereby providing an explanation of $\mathbf{E}_1$ by $\mathbf{a}_1$. I want this to be understood, not as “ $\mathbf{A}$ reducing $\mathbf{B}$ ”, but as “ $\mathbf{a}_1$ reducing $\mathbf{E}_1$ ”. Again, for this latter reduction to be symmetric it would have to be that we have “ $\mathbf{E}_1$ reducing $\mathbf{a}_1$ ”, thereby explaining $\mathbf{a}_1$ by $\mathbf{E}_1$. Once we are careful in distinguishing the *explanans* from the *explanandum*, cases of apparently symmetric explanations become clarified. Two reductions that at first past may seem aptly summarized by “biology to chemistry” and “chemistry to biology” – giving the appearance of symmetry – are found after further scrutiny to be better represented as “genetics to molecular chemistry” and “DNA mutation rates to fitness optimization constraints” – two reductions that collectively do not make a case for symmetry.

In §3.1, Nickles discussed how reduction₁ and reduction₂ each operate in different directions. First notice that α reducing₁ β and β reducing₂ α would not constitute an instance of symmetric reduction – as Nickles intended, the subscripts show that there are different notions of reduction operating. Instead we would need to have a symmetric reduction₁ or a symmetric reduction₂, a reduction that operates both ways with respect to the same concept of reduction. For a reduction₂, this would involve us performing “approximations of many kinds” from α to arrive at β , as well as using different approximations to make β approach α (Nickles 1973, 181). For the case of SR-CM, this seems unachievable. But the simple symmetrical cases provided earlier in §5.1.4 and §5.1.5 seem like they would be legitimate cases for a symmetric reduction₂.

At first glance, when the goals are either ontological or epistemic, it seems undesirable for one to create a reductive model that allows for symmetry. Yet we find that only Suppes and Schaffner have avoided doing so. To rectify this mismatch, we should look at each of the other authors in detail. Recall that symmetry was permitted in Kemeny and Oppenheim’s reductive model by being critical of the requirement that “ α is at least as well-systematized as β ”; specifically I focused on the lack of precision inherent in the concept of “systematization”. Sklar, rightfully, identified that the “systematizing” feature was introduced to prevent symmetric cases (1967, 114). Systematization was meant to be a quality that ordered theories: given two theories that concerned the same domain, one must possess a greater systematization than the other. Since I suggested that it could not successfully do this in some cases, symmetry was allowed in such instances. So despite Kemeny and Oppenheim’s model allowing for symmetry, I contest that

they would be displeased to find that this had occurred – and here because they would be displeased to find that “systematization” could not pull the weight intended.

In §5.1.1 we were able to produce a symmetric case with Nagel’s theory because, among other things, there are bridge principles that biconditionally link the exclusive vocabularies of α and β (refer to [5.1] and [5.2]). In response to an unrelated criticism (Kemeny and Oppenheim 1956, 10), we mentioned that Nagel considered amending this requirement in a footnote by changing the biconditionals in his bridge principles to conditionals (Nagel 1979, 355, ft 5). The suggestion would place a term β in the antecedent and terms of α in the consequent, (as made explicit in §1.1.1 and footnote 1.8). Were such a requirement in place, this would be sufficient to make the model antisymmetric.

The example provided of a symmetric Nagelian reduction warrants further examination. Without details of Nagel’s account other than that it required derivability, a fast analysis of a symmetric reduction might claim that $\alpha \vdash \beta$ and $\beta \vdash \alpha$. This might then lead one to claim that α and β were “identical” or “equivalent”, as was the case for Suppes in §5.1.3. But looking at the substance of the example, it should be clear that the constituents of each theory are not “identical”. For instance, one theory has $n+2$ predicates, the other has $n+4$. A better way of stating the relationship is both theories describe adequately true and false statements about the same range of phenomena. Anything that can be said with α can be said with β , and vice versa. A good analogy here would be language: I can describe adequately the contents of a room in Chinese or in English. Some descriptions are more verbose in one language or another, but the content is still the same. Any description in one language may have its content translated into another by a lexicon, likely one that relies on biconditional-like propositions. Here there is no “best language” in the absence of any criterion with which to judge it. Only when I want to privilege the description with the “least number of typed characters” or the “fewest number of syllables” will we be able to say which language is better – in regards to a specific situation/description or in general.

A symmetric reduction for the Nagelian model would operate the same way. One theory would provide a scientific account of the world, and we could just as well have used another. The descriptions provided, ultimately, have equivalent empirical predictions and isomorphic logical features. Preference of one over the other must be given only with regards to specific virtues. “Familiarity” may be one. For example, most people prefer to calculate by decimal

numbering rather than hexadecimal. Some frameworks may be more complicated for certain tasks, such as using only $\{\&, \sim\}$ when doing sentential logic. Regardless of the given criterion, the important lesson is that a symmetric reduction is not necessarily a strange notion. Instead it shows, for the Nagelian case, how two theories may be understood to be interchangeable.

This idea of interchangeability may be stretched to describe the symmetric cases from the New Wave model. In §5.1.5, I provided a case where we limited different aspects of α / β to arrive at β / α . This shows how, with regards to specific variables, each theory can be made to look like the other. In general, it shows us that either theory can be used in place of the other, given certain conditions that screen-off the relevant differences between the two. We may have confidence in one theory describing a situation, so long as this situation is one that accurately fits the limiting processes found in either reduction. In these cases one would have to be very aware of the $[A_n]$ and its limitations⁹⁴. There may even be *bona fide* explanations of α by β and β by α , so long as these occur in the appropriate contexts and respect the scope of each limiting process. This is not a case of different theory-parts explaining/being-explained-by other theory-parts, as we clarified above, but instead an instance where the same equations in each case swap roles as *explanans/explanandum*. A symmetric reduction would recognize that each theory could do better than the other in certain contexts. In some cases they are interchangeable, but in other cases, one or the other provides a clearly better explanation. However, in general, there need not be a “best theory”.

In most cases, a symmetric reduction would be confusing: either physics reduces to chemistry or chemistry reduces to physics, but certainly not both. In regards to most of our goals, symmetric reductions do not help in achieving them; indeed for many cases they instead obfuscate the goals. Two of our case studies could not be symmetric: the GR-NM case would not serve as an example of a symmetric reduction, nor would the SR-CM case. But the work of §5.1.1 and §5.1.5 show how a symmetric reduction might be achievable and, furthermore, quite reasonable. Such cases provide information about the relationship between the two theories, about when they are interchangeable, or how each may be said to explain each other.

⁹⁴ In relation to our quantitative example from §5.1.5, this would require that we keep the domain of validity, δ , in mind.

§5.2.3 Transitivity and Succession:

Fundamentalism in the sciences typically claims that there is ultimately one scientific theory that will be applicable in all cases, and all other potential scientific theories will either be incorrect or reducible to this final theory. For the physicalist⁹⁵, this is primarily an ontological claim. Micro-physics gets to decide what the ultimate constituents of the world are: atoms, fundamental particles, wave functions, strings, etc. As evidence for physicalism, one could attempt to show that chemical entities were “nothing more than” physical entities. This would perhaps be accomplished by a reduction from a theory that posits more-fundamental entities to one that involves less-fundamental entities, such as a reduction from chemical theory to physical theory. Here the details of the reduction would have to be robust enough to be justification for the ontological replacement. However, when seeking to reduce psychology to particle physics (or whatever other physical theory is to be considered basic), we would like to proceed through intermediate theories such as perhaps biology or chemistry. Without the stepping stones granted by transitivity, attempting to proceed directly – say from psychology to particle physics – is simply too big of a jump. Issues of complexity and grouping arise, blocking any hope of showing such a reduction regardless of confidence that such a maneuver should be at least possible in principle. Allowing for transitivity, or at the very least a large number of transitive cases, affords the physicalist intermediary branches that have potential to eventually connect the leaves of the tree to its roots (Oppenheim and Putnam 1958, 7).

Ontologically, one could grant that the particles of micro-physics were in no way more primary than the chairs and trees that we encounter every day. Still, one might think that there is good reason to privilege one physical theory over other scientific theories. One could believe that, more so than any other theory, physics provides the best explanations of the world. For example, chemical theory might do an excellent job describing why a specific reaction occurred the way it did, but there are facets of the theory – constants, equations, properties – that are simply put forth as “givens”. For example, that the different elements have different masses must be taken as an unexplained brute fact. Transitivity for explanatory reductions allows for one

⁹⁵ Some use “physicalism” as a thesis applicable to the philosophy of mind, asserting that there are no mental substances and that mental states are brains states. In this respect it is similar to materialism. I use “physicalism” here in the scientific unificationist sense, making it a (mainly ontological) statement about the sciences, as §G makes clear.

theory to be able to explain every other theory. Remember that epistemic reductions typically have theories explaining other theories, not theories explaining the world. Inasmuch as one wants a reductive model to allow for the possibility of there being one final theory that explains the nuances of every other legitimate scientific theory, then one will seek to include transitivity as a necessary feature.

Epistemic issues in successional reductions often emerge when looking at how theories relate temporally. As the name indicates, a successional reduction typically allows a later theory to explain its predecessor. If we allow that reduction may be transitive, then the current theory may be said to explain all past theories. This is desirable because it begins to paint a tacit picture of progress, as well as affirming successes by demonstrating how a current theory is superior to those in the past. As we demonstrated in §2.2.2, §3.2.1, and §3.3.1, sometimes the reduction's *explanandum* provides details about science's interface with history, such as "how successful a prior theory was". If our goal was also to transfer confidence to a new theory, as §3.3.2 established, transitivity is also a boon. It will allow epistemic weight to be transferred to the new theory from *all* prior theories.

As Nagel, Kemeny and Oppenheim, and Suppes each offer models of reduction that are transitive⁹⁶, I think that it is not unreasonable to infer that some of these philosophers might have been motivated by fundamentalist inclinations. Fundamentalists and physicalists need a reduction to be transitive if a reduction is going to be useful in establishing their worldview. Epistemically and ontologically, a transitive reductive model will be best able to support and be consistent with fundamentalist sympathies. But notice that Schaffner and the New Wave have models that are not always transitive. Rather, in each case the acceptability of a reduction depends on the strength of an analogy, and sometimes two tenuous analogies will not be enough to assure transitivity. This does not in any way squash the fundamentalist's reductive aspirations, but it does mean that each possible transitive case must be checked individually. Each of the first three models lacked the ability to account for approximations. As a result, they were each transitive in all cases. When an account of reduction allows that the reduced theory may be close to a reducing theory, the question becomes "how close?" Any time a quantitative

⁹⁶ Technically, Kemeny and Oppenheim's model is not transitive only in the cases where $\alpha = \gamma$, as we noticed in §5.1.2. This is merely a logical consequence of the fact that the model is never reflexive and there are some allowed symmetrical cases. Thus there are cases where $\alpha R \beta$ and $\beta R \alpha$, and transitively this would imply $\alpha R \alpha$, a contradiction with the impossibility of reflexive cases. I see these cases as aberrant, and, am content to claim that they are generally "transitive".

stipulation is given in answer, the intransitive cases of §5.1.4 and §5.1.5 can be created. The fundamentalist may still be able to arrive at their pyramid by virtue of the “step-skipping” that transitivity provides. Or perhaps there will be additional work with additional theories so that the analogies remain tight enough to give confidence in the chain. However this leaves the fundamentalist project, Schaffner and the New Wave disagree with authors such as Niebergall about the necessity of transitivity for models of reduction (Niebergall 2002, 148).

Wimsatt, thinking of accounts of reduction like the New Wave’s, claims that “successional reductions are intransitive” (1974, 677). Here he could be claiming two things: a transitive successional reduction cannot exist, or not all successional reductions are transitive. Because later he claims that “[successional reductions] would usually be intransitive”, I think that the latter is the most appropriate (and charitable) interpretation (Wimsatt 2006, 449). It is interesting that Wimsatt claims this only for successional reductions, for as §5.1.5 shows, inter-level reductions also are not transitive in all cases.

§5.2.4 Ranking Theories to Give Structure to Science

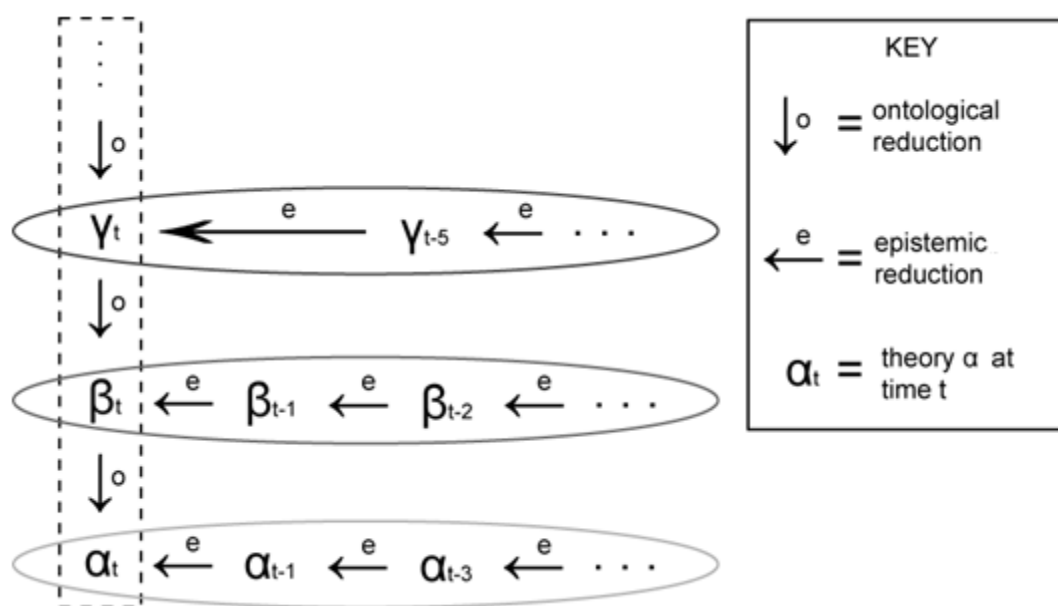
As mentioned before, antisymmetry and transitivity will *order* the set of scientific theories. This allows one to rank elements in terms of what may be reduced to what. Such an ability certainly interests the fundamentalist: if we allow that each current theory (that is the accepted physical theory, the accepted chemical theory, etc.) is comparable to each other (but not necessarily themselves⁹⁷), then we have a *totally-ordered* contemporary science that can be characterized by the scientific pyramid. Here every theory could be ranked in reference to another. There would be an “uppermost” theory, and there would be a “lowermost” fundamental theory that every other theory was reducible to.

We have excluded past theories in our total-ordering of “presently-acceptable theories”. If we instead focused on successional reductions, and have a reductive model that is antisymmetric and transitive, we might be able to capture a notion of scientific progress over time. If each prior theory was reducible to its successor, then we have an ordering which might

⁹⁷ Mathematical notions of “total ordering” or “partial ordering” typically require reflexivity (Weisstein 2015a) (Weisstein 2015b). As §5.2.1 dispensed with the need for reflexive reductions, I have appropriated these terms to apply yet lack the reflexivity clauses. In doing so, I think I have included the spirit of what these terms typically mean in a set-theoretic context without forcing reflexivity as a condition for reduction.

be said to demonstrate how science “gets closer to the truth”, granting some notion of approximate reduction. The conglomerate of each successional theory set might then be given as evidence against pictures of science that seem to eschew the notion of progress, such as Thomas Kuhn’s (1962).

These groups of ordered successor sets might also be employed to object to the pessimistic meta-induction argument. Traditionally construed as an argument against scientific realism (Laudan 1981), it begins by noting that prior scientific theories have made significant errors about how the world works and what it consists of. This then serves to undermine our confidence in the veracity of our currently-accepted theories. Current particle physics claims that there are quarks, yet it is likely that science errs here, just as when physicists claimed there was æther or phlogiston. Now the fundamentalist can point to a foundational theory that the entire body of scientific work will be reducible too, despite the likely mistakes of our current theories. We are “making progress”, and there is a kernel in each theory that has “gotten something correct”, even if just approximately.



{Figure 5.4}

Combining the inter-level and intra-level orderings, we are left with a *partial ordering* of science composed of several comparable subsets. There is an inter-level ordering across theory-

levels of the present time that preserves fundamentalism⁹⁸. There is also a successional-ordering that temporally reinforces a notion of scientific progress. In sum, from this perspective the scientific endeavor, across many disciplines and over its long history, has progression and coherence – it is unified. {Figure 5.4} aims to visualize this conception of science, labeling the inter-level reductions as “ontological” and the successional reductions as “epistemic”. This unification could then represent a structure for scientific endeavors that might preserve it against larger criticisms of irrationality and incoherence. For example, science could be viewed as rational because there was a progress across time facilitated by a notion of successional reduction that had “improved accuracy” as a function. Likewise, we might be able to understand science as coherent because inter-level connections bound the theories by reductions that showcase mereological dependencies.

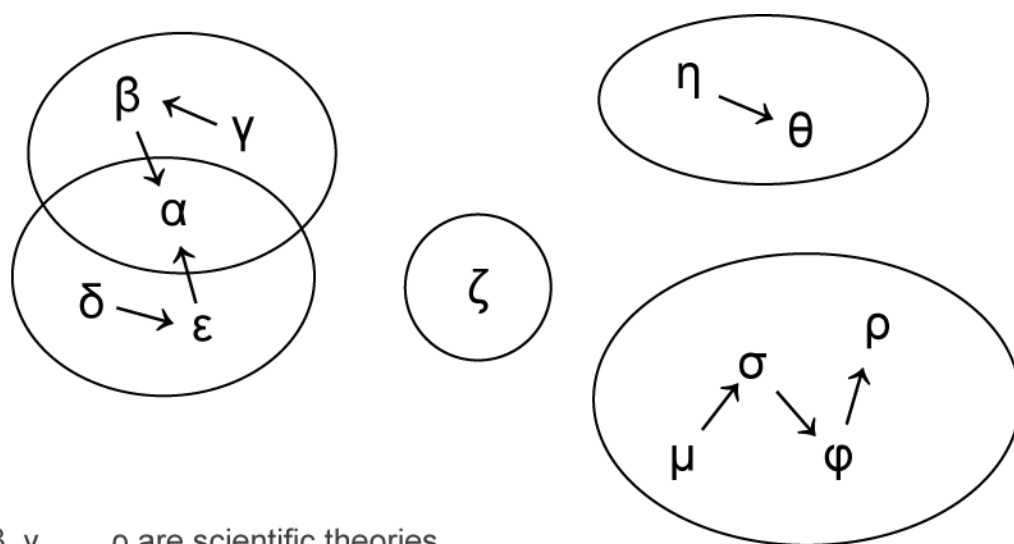
To further speculate, the unification provided by a partial ordering of the sort depicted in {Figure 5.4} might even be used to solve issues of demarcation: if a theory was comparable to some theory that was a part of the pyramid through the reductive relation, then it may be said to be a “scientific” theory⁹⁹. If we cannot reduce it or show that it reduces any of these accepted theories, we might then claim it to be genuinely different from what we have come to know as “science”. Notice that the notion of “science” employed here is historical, and not essential in any significant way.

In contrast to the “unity of science” partial ordering, there is another significant way that we might be able to structure theories in relation to each other with regards to a partial ordering. It might be that there are small groupings of comparable theories, perhaps across a brief period of time, perhaps across levels of organization, or perhaps both. This is scientific pluralism: a dappled world of internally-comparable theory groups. Within each group, there would still be a partial ordering and likely a sense of progression – concepts that some would hold as essential to any successful picture of science. To help visualize this picture, refer to {Figure 5.5}. However to speak of relations (that would be borne out by means of successful reduction) *between* each of the theory groups becomes tricky. In some ways, inter-group comparisons would be fruitless, as we would treat each theory-group as its own island; members of differing islands need not be

⁹⁸ We might even look for inter-level orderings that cut sections across past times. This would show that in past epochs of scientific development, there was still a reductive coherence.

⁹⁹ Such a move is very similar to the goal discussed in §3.4.1, where status was granted to past theories by virtue of their relation to other theories.

comparable. Such a picture is similar to Kuhnian incommensurability (Kuhn 1962). For others, there can be many theoretic relations between reduction-ordered groups, and these relations can be quite beneficial. For example, Alison Wylie demonstrates how archeologists may rely on disparate areas of science to provide several sources of evidential weight for claims, such as those employed in radio-carbon dating (2000, 233). Oftentimes evidence is derived from empirical claims in combination with theory. Were it the case that the theories employed when assessing different sources of data came from separate reduction-groups, then we could consider the results more robust epistemically. Errors in a single theory will not transit to the other evidence that relies on other theories, so having a broad theory base is in this way beneficial.



{Figure 5.5}

{Figure 5.5} is one way that we could make sense of the pluralist's worldview, where the arrows represent reductions. However many others exist that simply eschew the need for any ordering for scientific theories. It might be that any "communication" between theory groups facilitated by reduction simply should not be thought to be antisymmetric or transitive. There could still be theoretic "islands", yet the landscape of each island is not necessarily ordered in the mathematical sense. It may just be that different laws operate over different ranges of

phenomena, inter-level or intra-level, a picture quite similar to that described by Nancy Cartwright as a “patchwork of laws”¹⁰⁰ (1994).

To conclude, it might initially appear that fundamentalists have strong reasons to prefer a partially-ordered science, and thus have reason to embrace a model of reduction that is antisymmetric and transitive. However the possibility of creating a partial ordering in a different, pluralistic way shows that making a model of reduction with such features is necessary to exonerate the “unity of science” position, yet not sufficient. Furthermore to make all of science *totally ordered*, such that every scientific theory was comparable to any other, seems hopeless. To do so, we would have to allow successional reductions across time to be inter-level-comparable across layers of organization. To show by example: psychology might inter-level reduce to modern genetics; also Mendelian inheritance theory might successional reduce to modern genetics. However a totally-ordered science would require that either Mendelian inheritance theory reduce to psychology or psychology reduce to Mendelian inheritance theory – a bit of a stretch. This would place demands on a model of reduction that would force it to be trivial or misrepresentative in an attempt to accommodate a total ordering.

§5.3 Conclusion

This chapter investigated how reductions might be said to order theories in relation to one another. To do so, §5.1 examined five of the major models of reduction from earlier chapters to consider possibilities for reflexivity, antisymmetry, and transitivity. Reflexivity always occurred for the models of Nagel, Suppes, and the New Wave, while it was disallowed by the models of Kemeny and Oppenheim and Schaffner. Suppes and Schaffner had models that were always antisymmetric, whereas Nagel, Kemeny and Oppenheim, and the New Wave all provided models that allowed for symmetric cases. Finally the models of Nagel, Kemeny and Oppenheim, and Suppes were always transitive, while Schaffner and the New Wave models gave cases that could fail to be transitive.

§5.2 examined each of the three ordering relations with these results in mind. Reflexivity was deemed an undesirable and trivial relation for reduction. To show this, §5.2.1 looked at the

¹⁰⁰ Cartwright refers to the type of pluralism described as “metaphysical nomological pluralism” (Cartwright 1994) (Cartwright 1999).

goals that one might accomplish in a reduction and found reflexivity did nothing to support them. Additionally, the wordings of the five models could easily be changed so as to allow or disallow reflexivity quite easily, yet there was no univocal decision to prefer reflexivity or irreflexivity.

§5.2.2 vindicated antisymmetry as a desirable feature, on the grounds that it would preserve commonplace reductive goals such as successional progress and inter-level compositionality. The cases of symmetry that were possible need not be deviant cases, but instead can be ones that illustrate interchangeability. Transitivity was the subject of §5.2.3, a relation that is important for bridging long gaps of time for successional reductions, and large differences in scale for inter-level reductions. Only the two approximative models, Schaffner's and the New Wave's, prevented transitivity from holding in all cases. This is understandable given that in each case reduction hinged on the analogy made between the succeeded theory and the constructed analog theory – there are instances where the analogy is stretched too far.

§5.2.4 concluded by showing how the three relational properties may be employed to order science. The fundamentalists' hope was that an antisymmetric and transitive reduction model would lead to their desired scientific pyramid. Work here showed that such a result follows only if it is required that for every pair of theories, there was a reduction going in just one direction. This provision is excessive, as it requires theories distant across inter- and intra-level dimensions to nevertheless be somehow comparable. Once we allow that two theories need not have any reduction between them, it is possible to order theories within different clusters that are exclusive in their reductive communication. In this picture there is still a notion of localized order, yet there is not overarching order such as that in the scientific pyramid. Thus antisymmetry and transitivity are compatible with both the fundamentalist and pluralist worldviews.

Chapter 6: Functionalism and Goals

The perceived unity of reduction was an artifact of focus on structural or logical rather than functional features, when interests in reduction served foundationalist aims of increasing philosophical rigor, epistemological certainty, and ontological economy. These philosophical goals rarely matched those of scientists...

William Wimsatt, (2006, 448)

§6.0 Introduction:

This final chapter begins by summarizing the significance of the work done in the past five chapters. §6.1 lists and comments on the various goals recognized by examining the cases in the past chapters. Additionally the section highlights how the work of the past few chapters agrees and disagrees with other philosophers writing on intertheory comparisons. Notably, I recognize that my view on reduction is thoroughly functionalist, in the sense articulated by William Wimsatt. §6.1.1 comments on how intertheory relations and the case study from Batterman function in arguments for or against fundamentalism and pluralism. §6.1.2 considers the usefulness of a transitive reduction from Aristotelian dynamics to Newtonian mechanics to general relativity, an endeavor Wimsatt speculated to be fruitless. §6.2 provides a brief example of an inter-level comparison made by Eric Scerri, where he concludes that on the basis of recent quantum chemical calculations, chemistry is neither reduced, nor approximately reduced, to physics. In §6.2.1 I show how, by attending to the function intended for a reduction, there are some ways that this reduction could still be deemed a successful one. Finally, §6.3 recognizes the significant contributions that I believe this dissertation has made to philosophical considerations of reduction and intertheory relations.

§6.1 Reasons for Comparing Theories:

Scientists and philosophers are interested in comparing theories for a variety of reasons. Early literature focused on reduction, as we saw in §1. The first models of reduction could not

deal with approximation: each presumed that there was a complete agreement between theories over certain ranges. This is problematic because these models disallowed that either of the theories could make errors in their predictions of the world; the history of science shows that the dominant scientific theories have changed many times over the years, and the reasons for such changes have almost always been the result of the failures of the prior theories to give an adequate account of empirical phenomena. Thus examples of successional reductions rarely fit any of the models found in §1. Nickles, Schaffner, and the New Wave all wrote with the goal of allowing approximate models of reduction. The difficulty arises when attempting to adequately account for approximate reduction in a way that is not so permissive as to allow problematic cases, nor so stringent as to disbar genuine scientific candidates.

A clear lesson from §3 is that we should be wary of how to construe approximations. Batterman's interpretation of Nickles's reduction₂ was too narrow: by merely allowing a limiting relation between curves, we admit cases with ludicrous curves that nevertheless limit to Newtonian mechanics, and must ignore cases with curves that are extraordinarily close to one another yet do not have the property of equaling each other in the limit. The moral is that a limiting relationship on its own is uninformative. Instead, attentiveness to the empirical circumstances that areas of the curves inhabit allows for correct judgments. Indeed without these details, the payoffs that emerged in §3.2.1, §3.3.1, and §3.4.2 would be obscured. Batterman expresses a similar caution in his response to Belot in §4.4, where he warns of the errors that can occur when we presume reduction to merely be a problem of "pure mathematics".

One of the significant moves, beginning with Nickles, has been to distinguish between the different types of reduction involved: inter-level or successional/intra-level¹⁰¹. This is a distinction that has been maintained throughout this dissertation, however one which is often less prominent in the literature. Wimsatt has stressed this difference, claiming that "most accounts have conflated these two kinds of reduction" (2006, 448). The difference in some ways is methodological: inter-level reductions are often compositional affairs that seek covering explanations of phenomena, whereas intra-level reductions typically involve mathematical approximation techniques to show similarities in structures between two theories (Wimsatt 2006, 448–449). The other relevant difference between the two types of relations is one that is

¹⁰¹ Although I group them together here, not all intra-level relations are successional relations. For instance, I see the case study of §4 to be an intra-level relation that is not successional.

functional, a distinction that seeks to clarify the usage of reduction in light of the fact that inter-level and intra-level activities often proceed with different goals behind them. “*Inter-level* or mechanistic *reductive explanations* serve fundamentally different functions than *same-level* reductive *theory succession*, resulting in structural differences between them that had gone unnoticed” (Wimsatt 2006, 448, emphasis in original).

From the work of our case studies, I have demonstrated several goals that may underlie intra-level relations. We found that successional reductions may: **(i)** provide an explanation of (aspects of) the succeeded theory by the successor; **(ii)** provide an explanation of the theories’ successes/failures, as well as explaining details of the progress, both historical and conceptual, from the succeeded theory to the successor; **(iii)** transfer confidence to the successor from the succeeded theory; and **(iv)** delimit a range of applicability for the succeeded theory. I will briefly review the work that contributed to these findings.

(i) is a goal often associated only with inter-level reductions (Nickles 1973) (Wimsatt 2006). The work of §2.2.1 however demonstrated that it can also be a function of successional reductions. The fact that the inertial and gravitational masses are precisely equal to one another is a rather curious brute truth that Newtonian mechanics is unable to account for. The reduction of general relativity to Newtonian mechanics detailed in §2.2 shows how the two types of masses *must* be the same when Newtonian mechanics is obtained as a limiting case of general relativity (Weatherall 2011)¹⁰². Thus an important function of the reduction was to explain aspects of the succeeded theory, specifically, to explain the curious equality between the inertial and gravitational masses.

(ii) was present as a feature of an earlier model of reduction¹⁰³ (Schaffner 1977), and was recognized by other authors that focused on successional reductions (Sklar 1967) (Wimsatt 1974) (Wimsatt 2006). Notably none of the authors provided any of the details of how this goal might be realized, yet much of our work here showcases exactly how a reduction can accomplish this function. The limiting process of §2.2 served to “spread” the light cones at each point, causing the maximum possible speed of particles to increase. At the limit, there is no bound to the speed

¹⁰² Interestingly, part of Weatherall’s purpose in writing his 2011 paper was to demonstrate that how several predominant theories of explanation, such as the D-N model, were inadequate to fully describe the explanation provided by the reduction. This is quite similar to the goal that Batterman had in his early reduction work: to show how the explanations provided in intertheory relations are not easily captured by current models of explanation (2002b).

¹⁰³ Specifically I am referring to condition **(III)(ii)** of the Schaffner model: “ β^* explains why β was incorrect, and why it was once considered correct”. This was first presented in §2.1.

that particles may travel. Only at this limit are we able to begin to transform the spacetime into one that is classical. Detailed in §2.2.2, the GR-NM reduction explains what aspects Newton “got correct” or nearly correct (e.g. gravitational dependence on mass), and also explain where he erred (e.g. simultaneity and the “action at a distance” nature of gravitational attraction). §3.2, §3.3, and §3.4 concern the limiting of the special-relativistic momentum equation to a classical momentum. Empirically, the measure of momentum is very difficult to detect at typical earthbound velocities. Allowing the speed of light to increase without bound in the SR momentum equation, as is elaborated in §3.2.1, tells us how 17th century practitioners such as Newton might have been warranted in setting no upper bound to the speed of light. Another way of taking the limit is to allow the velocity to approach zero. Here we find a related description in §3.3.1 about the range of applicability of the theory, and about how long it took for the theory to be challenged by critical tests. §3.4.2 provides related functions under a slightly different framework (Rohrlich and Hardin 1983), showing a succinct way to specify the domain under which it would be impossible to distinguish the two theories, given the level of error in measurement.

Next, (iii) was the focus of §3.3.2, and noticed by (Nickles 1973) (Wimsatt 1974) (Rohrlich and Hardin 1983) and (Wimsatt 2006). Classical equations such as the momentum equation had been successful in explaining and predicting the world for a range of cases and over a long period of scientific scrutiny. By limiting the velocity to zero in the relativistic equation – and by keeping in mind the experimental capabilities during much of the period when the equation was presumed correct – we see warrant to transfer confidence from the successes of the classical equation to the special relativistic one. This warrant comes from the very close match of the two equations in the low velocity regime – a match that the reduction informs us of. Finally, (iv) saw mention in §3.4.1, and by (Nickles 1973) (Rohrlich and Hardin 1983) (Wimsatt 1974) and (Wimsatt 2006). The margin of error, δ , employed tells us precisely under what conditions we are allowed to employ the succeeded classical momentum equation without worry of error. If we want results to be within δ , then there exists a range under which the momentum equation will be applicable. Despite the known “incorrectness” of classical mechanics, there is still a large range of cases that the equations may be applied to.

§4.2 provided an intra-level case, one involving a “successor” and “succeeded” theory, that was not an instance of successional reduction. Batterman showed how there can be

intertheory activity involving asymptotics that requires help from the old theory to arrive at an explanation of universality. Is this an instance of reduction? A lack of reduction? Batterman is content to deem the affair an instance of an intra-theory “relation”. This is because the limiting involved was singular, and Batterman’s thesis was to claim that only regular limits could be candidates for reductions. However we observed that there are successional “physicist’s” reductions that are regarded as reductions despite the existence of singular limits – §2.2 provides such a case. I agree with Batterman that it is perhaps best to classify the case of §4.2 as an intra-level “relation”, but for different reasons. When explaining the rainbow, a covering explanation was facilitated by ray optics, because wave optics could not supply the necessary explanation on its own. This shows how, despite the admission that ray optics is in some sense “embedded within” wave optics, we should not dispense with ray optics too quickly. For the present time at least, it isn’t possible to get a wave-theoretic model of the rainbow off the ground in a way that adequately explains the universality of several rainbow features. Thus we aren’t trying to “reduce ray optics to wave optics”, “reduce wave optics to ray optics”, or both (in other words, the target is not a reduction₁, a reduction₂, or a symmetric reduction). There is no directionality in any significant sense attached to the affair, but merely two theories being cleverly employed to explain empirical phenomena.

How do the goals of intertheory relations that are not reductions differ from those of reductions proper? I see no great reason to think that the aims described by (i)-(iv) couldn’t likewise be employed in these cases. There was an explanation facilitated by §4.2, but not one theory explaining another, and not an explanation of progress. Instead, both theories provided a scientific explanation of the universality of a type of phenomena. In some cases this is an inter-level affair, when a phenomenon requires treatment from several different branches of science. For example, Eric Winsberg notices that nano-science often recourses to “multiscale” models, that are amalgams of theories from several different levels of the scientific pyramid, e.g. “quantum mechanics, molecular dynamics, and continuum mechanics” (2006, 584). In these cases an explanation is provided by an interface of many theories, but this clearly need not constitute an instance of reduction. Instead it is just an admission that sometimes a single theory is insufficient to deal with the phenomena. Interestingly, the intra-level case from §4.2 is evidence of a case representing the current indispensability of two theories that are both attempting to characterize the same phenomena in different ways.

§6.1.1 Pluralism and Fundamentalism:

Although it is never a term employed by Batterman, I think the case from §4.2 provides an example in favor of *scientific pluralism*. Batterman sought to demonstrate that to achieve an adequate explanation of a rainbow, it is necessary to draw on the resources of two theories: wave optics and ray optics. Yet it is unclear how we are to interpret “necessary” in relation to the metaphysical structure of science. The fundamentalist, championing the reality of one basic theory, would likely read Batterman’s conclusions differently from the pluralist, who rejects the need for a singular ultimate theory that has universal import in all scenarios. The pluralist could examine the wave/ray optics case and not see conflict with their worldview. The possibility of wave optics being unable in principle to provide an adequate explanation of rainbow phenomena would merely be an example showcasing the lack of a “final theory”.

The fundamentalist would contest that our current understanding of wave optics must in some way be deficient, and that more work will unveil the “overarching theories” necessary to provide an explanation of rainbows from wave theory *ab initio*. In this way the “necessity” of each theory to explain the universality of rainbow phenomena may be a result of the current theories’ distance from the final theory: right now, we must employ both. The necessity to employ of both theories is contingent, given our current understanding, but in the future when more physical (and likely mathematical) work is done we will be able to cease our reliance on past, erroneous, theories. Although not logically disallowed by the evidence provided, to maintain an entrenched belief in the truth of fundamentalism may seem stubborn when we consider how few challenges have been made against the viability of wave optics. Even if the fundamentalist is unwilling to admit that Batterman is correct in showing the indispensability of either theory for the explanation of rainbow effects, as Belot and Redhead do in §4.3 and §4.5, other cases exist where succeeded theories seem to do a better job describing certain phenomena, such as the success of demonstrating macroscopic chaos by classical dynamics, as opposed to modeling by quantum systems (Belot 2000). Collectively these cases may provide inductive ground to question fundamentalism.

The example of §4.2 thus provides evidence to dispute the fundamentalist project. This is ironic considering that the fundamentalist often seeks succor from successful reductions, as we saw in §5.2. If there is an accepted theory in a macroscopic domain, the fundamentalist will seek to produce an inter-level reduction to the microscopic fundamental theory. In this way the primacy of the fundamental theory need not be undermined by the existence of other successful theories from different branches of science. However the fundamentalist's claim is one that universally demands that every other accepted theory be reducible to the fundamental one. Thus existential cases of inter-level reductions should do little to solidify the existence of any foundational pyramids. This is because the existence of just one accepted and irreducible theory that does work explaining something which the fundamental theory cannot is problematic for the fundamentalist. So, in tandem with demonstrations on case-by-case bases of how each reduction obtained between the fundamental theory and all other tertiary theories, a preliminary move would be to champion a model of reduction that allowed for partial-ordering by antisymmetry and transitivity. But we know that this will not suffice, as §5.2.4 demonstrates that a partial-ordering does not necessarily commit one to a fundamentalist worldview. Instead the pluralist may have a dappled world that is nevertheless partially-ordered by relations from the very same reductive model. Furthermore, the existence of cases of symmetry and intransitivity in a model are problematic, as an endorsement of such models leaves open possibilities for structurings of science that are incompatible with the fundamentalist's *Weltanschauung*.

§6.1.2 Aristotelian Dynamics and Newtonian Mechanics:

Just as the functional approach may be used to judge whether a *reduction* has occurred, it may also be used to judge the adequacy of a *formal model of reduction*. There is a worry, first articulated in §5.1, concerning the grounding of any argument that might seek to endorse one model of reduction over another. Certain models, so long as they are consistent, may see little instantiation in genuine scientific cases, but this does not show that they are bad models. They could be classified as “rarities” or “empirically underdeveloped cases” – none of which speaks directly to their veracity. Instead the move made in §5 was to recognize that any feature a model

might possess, logical or otherwise, can be scrutinized in reference to the goals one would seek to accomplish by a reduction. We saw in §5.2.1 that reflexivity, endorsed by Niebergall, ultimately proved to be an uninteresting property for a model to either possess or lack. The reasoning behind our disagreement with Niebergall centered around the function served by a reduction. The goals of a reduction may play some role in the assessment of the philosophical activities surrounding intertheory relations.

Wimsatt makes a similar functionality-driven point when discussing the role of transitivity in successional models. He claims that a hallmark of a successional reduction is that it “relate the new theory to its *immediate* predecessor” (Wimsatt 2006, 449, emphasis in original). The reason for this is that none of the functions of an intra-level reduction can be realized by relations to long-distant succeeded theories. For example, function (iii) of “transferring confidence” would seem to be misplaced, for there is little confidence to garner from an antiquated theory in the first place. Wimsatt further quips “whatever the claims of historians, no scientist sees *scientific* relevance in tracing special relativity back to Aristotelian physics!” (2006, 449, emphasis in original).

This is an interesting example. Recently, Carlo Rovelli has demonstrated a reduction from Aristotelian dynamics (AD) to Newtonian mechanics (NM) (2013). In it, the author is quite explicit about the benefits: “the comparison sheds light on the way theories are related” (Rovelli 2013, 2). By some excellent exegesis of Aristotle, as well as some clever mathematical finesse, Rovelli is able to show in what ways AD may be seen as an approximation of NM. By manipulating the equations of NM – taking limits, eliminating square roots, etc. – these equations will reduce₂ to the equations that Rovelli extracted from Aristotelian canon. The reduction sheds light on how Galileo, Newton, and others were influenced by AD in the development of NM. Next, it sheds light on what influenced Aristotle in his development of his physics, such as how his domain was restricted to “*objects in a spherically symmetric gravitational field (that of the Earth) immersed in a fluid (air or water)*” (Rovelli 2013, 4, emphasis in original), and furthermore how Aristotle was able to provide an empirically accurate theory despite having limited mathematical tools and experimental evidence. Finally, the reduction sheds light on how Aristotle “worked to yield a highly non trivial, but correct empirical approximation to the actual physical behavior of objects in motion in the circumscribed terrestrial domain for which the theory was created” (Rovelli 2013, 7).

My goal is not to tout the laurels of another successional reduction; thus I have been brief in discussing the AD-NM case. Instead I would like to show how this case has the potential to give us insight into the benefits of transitivity for successional reductions. The “preposterous” transitive case suggested by Wimsatt would involve us beginning with GR, and attempting to show how we might manipulate the GR equations to arrive at AD. In §5.2.4, I provided a toy example of a transitive New Wave reduction. However, this case gives us guidance for how a *bone fide* example could proceed. Rather than reinventing the wheel, we may take the approximations and limits made on GR in §2.2 as our starting point. This will allow us to arrive at NM, and then from there we may proceed via the approximations devised by Rovelli to reach AD. This would be a real-world case of a genuine transitive reduction. Although I have done no more than sketch this reductive process¹⁰⁴, were we to have the details, would we be able to claim any scientific relevance from the result?

First, we would be able to delineate a domain of validity, and realize goal (iv). This would serve the historically helpful purpose of showing where and for what ranges AD applies. AD doesn’t have any explicit equations (during the Rovelli reduction the equations that were used had been extrapolated from the Aristotelian corpus), and I think in some ways that is a benefit. Few people employ *the equations* of NM in their everyday lives, but instead rely on some of the general, qualitative claims. In the same way, some of the claims of AD may likewise apply. One example is the extracted AD postulate stating: “heavier objects fall faster: their natural motion downwards happens faster” (Rovelli 2013, 3). Part of Rovelli’s project was to show how, when in a fluid, claims like these are empirically adequate. Thus operating in everyday life, Aristotelian tenants are not outlandish. Indeed after reading Aristotle’s *Physics*, when watching a few objects fall off my desk, I may remind myself of the above postulate to describe the motion – rendering Aristotelian physics actually helpful.

Indeed to show any usefulness of long-ago succeeded theories, the only relevant comparison would be a reduction to the currently-accepted theory. Take current theory α , succeeded theory β , and long-ago succeeded theory γ . A γ - β reduction serves to demark the domain of validity of γ in relation to β . But, because β is no longer the accepted theory, this does not tell us where we may be said to use γ – because we do not know under what circumstances we should be confident to use β ! To do so we must further consider where employing β is

¹⁰⁴This is a current interest of what work will follow from this dissertation.

warranted, by providing a β - α reduction. To sketch this for our case, we know from the AD-NM reduction that the AD claim “heavier objects fall faster” is valid so long as the objects move in a fluid, and also so long as they do not fall long enough to reach their terminal velocities. But we would be wrong to think that this held for all possible velocities. The NM-GR case would impose further restrictions, such as constraining the AD maxim to cases of “low velocities”, as we learned from §3¹⁰⁵. To stop the reduction chain short of the presently-accepted theory would potentially leave out the next iteration of scientific advances, thereby providing an uninformed “domain” for any applications of long-ago succeeded theories.

Another important benefit of an AD-GR reduction would be to learn contextual details about how AD and GR relate, as mentioned by goal (ii). As the pathway pursued in the transitive reduction proceeds by way of NM, it may likely be that the considerations gleaned from the AD-GR reduction add no further benefits beyond those provided by a consideration of the NM-GR reduction in tandem with the AD-NM reduction. *But this is still beneficial*. We may understand what concepts of AD “survived the trip” to GR. The reduction can tell us what Aristotle “got right”, where he erred and by how much, and finally help to explain why Aristotle made the choices that he did. Rovelli stresses how, after the reduction “one can still recognize old Aristotle's vessel, after quite many repairs and improvements, in the conceptual structure of modern theoretical physics” (Rovelli, n.d., 9). Seeing these details in relation to the choices made by GR certainly seems interesting to a scientist in the same way that the NM-GR reduction exhibited a similar payoff as discussed in §2.2.2.

Lastly, Rovelli’s reduction shows us that not all successional reductions need to involve the currently-accepted scientific theory to be worth pursuing. I do not think that the insights gained by the AD-NM reduction are necessarily *all* historical, in a narrow sense. This discussion has engaged with Wimsatt and also provided an actual example of a successional transitive reduction, giving it some voice in a discussion concerning the nature of theory progression and theory change. If we are to concede that part of the interest in reductions, scientific progress, scientific structure, and the rest is philosophical in kind, then this is enough to assure interest in an AD-GR reduction.

¹⁰⁵ Although §3 focused on the momentum equation of SR, a special case of GR, its lessons still apply to the domain of validity as it concerns our purposes.

§6.2 Scerri's Assessment of Quantum Chemistry:

To show one final example of how the functional view of reduction may aid philosophers writing on the subject, I include a brief case from Eric Scerri. Scerri's paper, "Has Chemistry Been at Least Approximately Reduced to Quantum Mechanics?", begins by giving a brief history of reduction in the philosophical literature. He observes that Nagel's derivational model is overly-constricting, and instead attempts to leverage a concept of "approximate reduction", one inspired by Putnam. Here is the Putnam quote that Scerri presents in the article:

It is perfectly clear what it means to say that a theory is approximately true, as it is clear what it means to say that an equation is approximately correct: it means that the relationships postulated by the theory hold not exactly, but with a certain specifiable degree of error. (Putnam 1965, 206–207)

Juxtapose this "perfect clarity" with the nuances of approximate reduction models introduced in §2.1, and we are left with a worry that Putnam's notions of approximate *truth* may change when being applied to approximate *reduction*. Putnam was presumably thinking about how we might go about judging the truth of a theoretical claim such as "neutrino flux is $1.3(1 \pm 0.6) \times 10^7$ neutrinos $\text{cm}^{-2} \text{sec}^{-1}$ at the Earth's surface" (Bahcall and Shaviv 1968, 113). Here the "approximate truth" of the statement could be made by establishing error in this *claim* in relation to reputable observations. However it is unclear how we are to proceed when verifying the details of a reduction of one *theory* to another *theory*. Despite these reservations, Scerri's goal is not to be specific about the details of degree-of-error reduction and how it might operate in general. Instead, he intends to show that any modest construal of an approximate reduction – details of the functioning of a "degree of error" notwithstanding – will answer his paper's title question in the negative once we examine the specifics of the two quantum chemical cases that Scerri presents.

Let's begin by sketching Scerri's two cases. Each result is from contemporary, computational quantum chemistry: scientists are attempting to predict chemical values by employing *ab initio*¹⁰⁶ quantum mechanical calculations. The first case examines the structure of

¹⁰⁶ In this context, an *ab initio* calculation is one which proceeds from a modeled wave-function of the atom and does not rely on any "help" from experimental information along the way. Naturally-measured physical constants are allowed, such as Planck's constant or the speed of light (Scerri 1994, footnote 1).

the CH₂ molecule to see if it is bent or not. Spectroscopic analysis, an experimental endeavor, initially concluded that the molecule's shape was linear. Researchers, proceeding *ab initio*, calculated that the molecule should be bent by 135.1°. This theoretic result led experimentalists to go back to reexamine the actual shape of the molecule. Upon doing so, the experimenters conceded that they had originally erred, concluding along with the theoreticians that the molecule was indeed bent. Such a result is thus cited as being evidence for quantum chemistry "coming of age" (Goddard 1985, 922).

Scerri is critical of the result because of the calculation methods employed. To arrive at the H-C-H bonding angle, complex equations derived from wave-function configurations must be solved. Here the method employed represents the wave function as an infinite number of one-electron functions in a series (Scerri 1994, 165). Computers and approximation methods are employed, as an analytic solution is not possible. As with any complicated infinite sum, there is a worry of convergence. Scerri points out that when approximating, in some cases there is no proof of convergence for the computational methods employed. Thus, when iterating an approximation technique for the series, stopping at a specific value gives no confidence that the arrived value is reliable.

In some cases there are proofs of convergence. However in these cases there are instances where the approximation methods employed give results that are not bounded when the results are given for a specified number of iterations. Without both an upper and lower bound on the values, we lack "error bars" and thus the ability to tell, even approximately, where the final value might lie. For example, if a supercomputer arrives at a value for the series by performing a billion iterations of a given solution technique, it is unclear how far away the number arrived at will be away from the actual solution to the series. The series may converge, and the approximate value will be somewhere along the path towards convergence, but one cannot know exactly – or even approximately – where.

Much of the creative scientific work involves researchers choosing which of the approximation techniques to employ, and which composition for the infinite series is preferable. Furthermore, there is the difficult task of choosing which of the numerous (on the order of billions) initial conditions of the molecular setup to employ when shaping the wave function itself. Scerri is skeptical about how the experimenters go about making these choices: when you have nothing but "a mixture of intuition and past experience" to guide you, the worry is that the

results will be unavoidably biased and the decisions *ad hoc*, even if inadvertently so (1994, 164). Thus it appears that Scerri's goal is to give evidence for why the result should be considered less than robust.

The other case which Scerri examines involves theoretic and experimental work to determine the symmetric stretching frequency of Si-C as observed in the Si₂C molecule. Experimentally the bar was set at 658 cm⁻¹ (Kafafi et al. 1983). Two years later a paper was published that gave a theoretic prediction of 823 cm⁻¹ (Grev and Schaefer 1985). Finally, there was a calculation done by experimenters in 1991 which set the line at 840 cm⁻¹ (Presilla-Marquez and Graham 1991). Later, two leading quantum chemistry groups brought their theoretical techniques to bear on the problem. The two groups ultimately disagreed; a group headed by Schaefer put the line close to 840 cm⁻¹, whereas another group headed by Handy found the line to be very close to 658 cm⁻¹.¹⁰⁷ In addition to this disagreement, Scerri notes how both groups had a notable difference in their results, as each chose to allow the computation process to run for different iteration lengths¹⁰⁸. Thus Scerri concludes that this case “does not say much for the reliability of current quantum chemistry, the claim that ‘quantum chemistry has come of age’, or indeed the claim that chemistry has been reduced to quantum mechanics” (1994, 168).

Before detailing my main problem with the Scerri paper, there are some smaller worries that are worth mentioning. Argumentatively, one might seek to fault Scerri for providing little discussion regarding what “approximate reduction” entails. However, I think that Scerri has tersely provided a reasonably clear picture. Here is my reconstruction:

If β *approximately reduces to* α within a “degree of error”, σ , then: the prediction for a value made by α and β differs by no more than σ .¹⁰⁹

This provides only one necessary conditions for approximate reduction¹¹⁰. However, as Scerri's argument only concerns the above essential feature of an approximate reduction, I will stop here.

¹⁰⁷ I rely on Scerri as a source for each of these points (1994, 166–167).

¹⁰⁸ I have omitted the details which Scerri's description hinges on (1994, 166–168). I do not see his summarizing assessment of the details to be particularly contentious. Finally, to an untrained practitioner, I submit that the material provided in his paper will likely be uninformative.

¹⁰⁹ Because I think that the measure of “degree of error” is similar to the δ of §3.4, I have employed σ to keep the two notions conceptually distinct.

I have kept Putnam's locution of a "degree of error", but it is worth mentioning how such a term might be misleading. Often in science the referent of "error" is an estimate of experimental deviance, perhaps resulting from inaccuracies resulting from equipment or experimenters. Here the term is to be understood in a mathematical sense, as a measure of how the two theories differ in their respective predictions.

I have been purposely vague about how we are to calculate σ , but I think that the details are slightly more complicated. Suppose, with regards to a quantitative prediction, α predicts **A** and β predicts **B**. Let α to be "quantum mechanics" and β to be "chemistry". To use $|\mathbf{A}-\mathbf{B}|$ as our means of calculating σ would seemingly be problematic as σ could be in any variety of units, including nanometers, joules, or ergs, depending on the nature of the thing predicted. A solution to this problem might be to speak of the difference between the two theories as a percentage: here troubles resulting from differing units and scales would seemingly dissipate. But how are we to calculate σ as a percentage? There are two options:

$$\text{The error of } \mathbf{B} \text{ in } \mathbf{A}: |\mathbf{A}-\mathbf{B}| / \mathbf{A} \quad [6.1]$$

$$\text{The error of } \mathbf{A} \text{ in } \mathbf{B}: |\mathbf{A}-\mathbf{B}| / \mathbf{B} \quad [6.2]$$

Typically, we are to divide by the "established value". The question of using [6.1] or [6.2] reduces to a question of which, either quantum mechanical value or the chemical value, we believe to be established. Regardless of whether we use [6.1], [6.2], or some other method that does not employ percentages, information concerning the difference may still be extracted. However I think that there are arguments in favor of choosing either theory to be the established one. One suggestion would be to always take α for our "primary" or "fundamental" theory, and thus rely on [6.1] as a measure of σ . Typically when doing routine science education laboratory experiments, we presume that the theoretical value is "established" and calculate the error of our measurements from the "known" value. But this is not our case. Both experimenters and theorists are competing for primacy: one of the major points of interest in the Si_2C case is that there is no clear "established" value.

¹¹⁰ There must surely be other requirements. Likely the range of predictions for each theory is not exclusive, as there will be phenomena that α can make that β cannot. How these are to be accounted for within a theory of approximate reduction needs to be stipulated.

Tacitly, Scerri has weighed in on the matter. In the article he makes an error calculation, stating that there is an error in the “chemistry” value (658 cm^{-1}) vs. the “quantum mechanics” value (823 cm^{-1}) of 25% (Scerri 1994, 166). Thus we see that he chose [6.2], as [6.1] would yield an error of 20%. This makes the “established” value to be the one that is employed by the chemical experimentalists. This should come as little surprise, given his reservations about the theorist’s methodology.

The above notion of “approximate reduction” makes reference to σ , yet Scerri does so only implicitly. His conclusion is clearly that we do not have an approximate reduction of chemistry to quantum mechanics. For him to arrive at this conclusion, it must presumably be that the σ observed in each of his two cases is larger than desired. But what is an acceptable σ ? How is one to be certain that the observed σ ’s are inadequate? The details for how we are to make these decisions for the general case seem difficult to provide. Presumably Scerri is relying on his experience as a chemist and as a philosopher when he passes judgment for the two quantum chemistry cases, but a worry remains that his decision is *ad hoc*. Even if Scerri’s rejection of the two observed σ ’s was the correct call, the relevant information is what seems most interesting. Scerri cites several features that informed his judgment about the viability of the reduction. The specific concerns include:

(I) The addition of **f** orbitals on the silicon and carbon atoms, which usually improves agreement with experiment in these types of calculations, produces a worsening in the frequency error in three separate methods (SCF, CI and EXT SCF), although the energy shows improvement.

(II) None of the [employed approximation] methods emerges as the clear winner in calculating fundamental modes from first principles. The outcome seems to depend on which particular mode is being considered.

(III) Overall, the error in ω_1 strays considerably from one method to the next and even on going to more extended sets within the same method. (Scerri 1994, 167)

Considerations such as these are likely what Scerri, *qua* chemist, supplies as evidence concerning why he is dissatisfied with σ . Importantly, it is not sufficient to determine the magnitude of σ and then make a judgment. Instead there is a complicated, nuanced story about

how such a number was generated to inform the judgment of those seeking to make claims about the adequacy of σ .

One additional worry concerns the actual nature of the relation for this case. Scerri himself has billed this as a “chemistry to quantum physics” reduction, seemingly putting “chemistry” as our β and “quantum mechanics” as our α . But are these both theories? Certainly our *ab initio* quantum calculation follows from the dictates of “quantum mechanics” in the fullest theoretical sense. But are we comparing it to chemical *theory*? Instead, it would appear that we have spectroscopists who are simply measuring a value, albeit via a very difficult and nuanced process. Even if we grant that “chemists” are performing the measurements, how entrenched is modern chemical theory in the determination of the experimental results? It would seem that there is an argument that suggests that it is merely an empirical result that is being presented, not one that is a necessary constituent of “chemical theory”. In this respect we would actually be carrying out a test of how well quantum mechanics (coupled with modern approximation techniques and supercomputers) predicts an empirically measured value such as a bond angle – a datum that has a common chemical usage. The worry is that we are left with not a “reduction” of one theory to another theory, but instead a quasi-“confirmation” of sorts.

§6.2.1 The Functionalist Therapy:

The work of the last few chapters has showcased the multifarious purposes that we might have when comparing two theories. While each of our case studies was intra-level, I believe that they provide insights that may also be helpful for inter-level cases such as this one. Scerri was critical of the lack of precision and accuracy displayed by the theoretical calculations. In addition, he gave the details of his reservations concerning the methodology of the theorists, as we saw by (I)-(III) above. Scerri was able to conclude that the σ was simply too large in relation to considerations such as (I)-(III). But what was the purpose he had in mind for the inter-level comparison? Surely different purposes warrant different considerations, and for some functions observed σ 's such as the ones discussed in his two examples might indeed be within the bounds needed to satisfy the requirements of an “approximate reduction”.

Scerri desired that quantum chemistry be able to predict the experimental values much more reliably than the theorists currently have done. The underlying hope of this desire was presumably centered on the ability for quantum chemistry to serve as a predictor of chemical properties, to a degree of error that would not significantly affect those doing chemical experiments, calculations, and theory. At the time of his article, Scerri recognized that quantum-chemical predictions were far away from the values provided by experimentalists – simply too far away for most chemical purposes. Thus we could not rely upon quantum chemistry as a replacement for experimentally-obtained values, on the grounds that quantum chemistry simply does a poor job of adequately capturing the empirical.

We cannot dispense with chemical experimental evidence and rely heavily instead on quantum chemical theory to make predictions and judgments. This seems to be the sort of claim made in rebuttal to anyone who would seek to claim quantum chemistry as a *replacement* to experimental chemistry, in the sense that the calculations and predictions quantum chemistry provided were as reliable and as robust as the experimental values. The values that the two methods provide are not really that “close”. This is similar to the judgments that could be made of the analogies $[A_n]$ and $[A_s]$, in §2.1 and §2.1.1, especially in the quantitative sense that they were construed in §5.1.4 and §5.1.5¹¹¹. Harkening back to §1.2.1, it seems that the associated goals are perhaps ontological in nature: “can the stuff of quantum mechanics predict the stuff of chemistry”? Scerri’s answer is “no”, or minimally “not very well”.

But what are the other goals that we could put the comparisons to? The article by William Goddard, “Theoretical Chemistry Comes Alive: Full Partner with Experiment”, which provided the “coming of age” quote cited by Scerri, provides a nice alternative. Goddard’s purpose, reliably indicated by the title, is to show that theoretical quantum chemistry is now capable of making predictions that may be taken to have similar epistemic weight to those produced by experimentalists. The article examines one of the examples that Scerri was interested in: the bending of the CH_2 molecule. The absence of error bars were Scerri’s primary concern in this case. Goddard seems less interested in the magnitude of error bars and more focused on the successfulness of the challenge presented by theoreticians to the experimental consensus of the molecule’s linearity. Earlier in the article, Goddard tells of a time when “one could not trust

¹¹¹ Here we employed δ as a measure of the quantitative “distance” between α and β , but when considering Scerri’s “approximate reduction”, we could instead make a similar case for σ .

theory (except when it was applied to small molecules such as H_2) unless it was confirmed by experiment, and theorists would generally have been foolhardy to suggest that theory was correct when there was disagreement with experiment” (1985, 917). To exhibit the shift, he merely needs to supply some historical cases that exhibited these qualities; in this regard his paper does a fine job. Goddard does not presume to claim that theory has supplanted experiment, or that in the majority of cases we can rely solely on theoretical results. Instead he is trying to convey his excitement resulting from the changes in the discipline, and as such can be read as having the following goal: to show that theoretical quantum chemistry has the legitimacy to inform, confirm, and collaborate with experiment.

There is another goal that we might read into Goddard’s activity, and into the overall function we might have for a reduction between theoretical quantum chemistry and experimental chemistry. In §3.1.4, Rorlich and Hardin sought to show how a reduction could vindicate the intra-theoretic progress of science over time, a view that was also discussed in §5.2. One might envision an analogous ontological progression occurring in the quantum chemistry case. At first there is the lower-level theory, quantum chemistry, in its infancy and unable to reliably engage with the claims, experimental or theoretical, of the upper-level science of molecular chemistry. The mereological complexities, as well as the fact that quantum chemistry is a newer field, result in the lower-level theory simply taking some time to “catch up”. Another example of this would be the slow development of the chemical underpinnings of inheritance: for a time, chemistry could not account for the predictions made by rudimentary genetics. But just as scientists were able to gradually articulate the molecular mechanisms of inheritance, so too do we see quantum chemistry slowly making these steps (Hershey and Chase 1952). The relevant goal here would be “to showcase the progress of inter-level science, by having lower-level theories be able to predict and explain mechanisms operating in upper-level disciplines”. Even if this function is not one that Goddard explicitly had in mind, it nevertheless is a goal that can be attributed to the above activity.

In the sense of solidifying ontological primacy, or assuring the dispensability of chemistry, I think that Scerri has done an adequate job in deciding that we should not consider chemistry “approximately reduced” to quantum mechanics. Scerri read each of the quantum chemistry papers and then looked to see how well the details served as evidence for a reduction (approximate or no) from chemistry to quantum mechanics. However reductions, “approximate”

or otherwise, may service many functions. By examining the claims of other scientists discussing the results, different goals emerge. Here the inter-level relation may exhibit the progress of quantum chemistry, and show how it may now be a more serious contributor to problems in contemporary chemistry.

§6.3 Contributions and Conclusions:

This dissertation has sought to show, through examples, the benefits gained by an attention to the many goals that underlie intertheory relations. My contribution to the existing literature is to show the details of how goals can be realized by the many types of reductions, and how various types of reductions can be understood to realize certain goals. At one time or another, other philosophers have mentioned these functions, but in notable cases they have failed to show the details in relation to scientific examples.

Hooker suggested that the GR-NM case was a good fit to the New Wave model (1981b, 203), and the work of §2 vindicates his intuition for the first time. §3 discusses the SR-CM reduction, an example of critical importance to the successional reduction literature. After providing a discussion of every major author to discuss the details of the case, each has failed to recognize the significance of the context surrounding the limiting process. §4 provides a complete summary of the debate surrounding one of the central cases in Batterman's recent book, and further grounds the discussion in relation to other historically relevant conversations about reduction and theory comparison. Additionally this section shows the falsity of one of the main theses of *The Devil in the Details*, the claim that cases that hinge on singular limits cannot be instances of reductions (Batterman 2002b). §5 provides logical analyses of three ordering relations for the five most influential formal models of reduction in the literature. Surprisingly, some models allow for cases of symmetry, a result that hasn't drawn the attention of authors examining, or even propounding, those models. A functional analysis is applied to these ordering relations, concluding that Niebergall's desire to preserve reflexivity was unfounded (2002). The chapter further shows how the ordering relations might be said to structure scientific theories, demonstrating that even a partially-ordered reduction does not necessitate the fundamentalist pyramidal worldview. §6 concludes the discussion by relating my work to those of others in the

field. It notices the potential to view Batterman's work as a case for pluralism, and shows, *contra* Wimsatt, the promise of examining transitivity in the AD-NM-GR case. The chapter finally provides an example of how the work of the past 5 chapters may be employed in an analysis of a contemporary philosophical paper concerning reduction, Scerri's "Has Chemistry Been at Least Approximately Reduced to Quantum Mechanics?"; by reexamining the goals involved in the reduction, there are ways in which we might answer "yes", despite Scerri answering "no".

To conclude, I am quite pleased with my contribution to the ongoing intertheory literature. This dissertation doesn't propose a radical new model of reduction, or claim to have conceived of a new way of viewing intertheory activity that will turn the discipline on its head. My work, taken as a whole, cannot be said to univocally agree or disagree with the majority of other writers and viewpoints – in most cases I take issue with smaller details and their articulation. I add many case-driven insights and provide examples, addendums, and revisions to the large and small points previously recognized by other thinkers. None of this should be taken as a failure of this dissertation; instead, from a long history of many informed, intelligent philosophers of science all carefully writing on a delicate subject, I look back on the work of this dissertation and believe that this is exactly what success and progress consist in.

§G Glossary of Important Terms

Preliminary Note on Symbolization: In this document, I unify the symbols employed by myself and other authors to make this easier for the reader. Oftentimes authors will use differing syntax, subscripts, etc., and I believe that by standardization there will be less cross-referencing and annoyance to the reader. References to scientific theories generally are done as α , β , γ , etc. Thus a quote that may have been originally stated as “theory X_1 to an earlier theory X_2 ” will be here rewritten as “theory β to an earlier theory α ”. For reductions, in the relation “reducing β to α ”, I will most times refer to α as the “reducing theory” and β as the “reduced theory”. Any usage of letters *qua* symbols will be done in bold.

Bridge Principles- A set of statements defining the unshared vocabulary of one theory in terms of the other theories vocabulary. Logically, these are often biconditionals. *Example:* Defining “neutron” from atomic physics in terms of “quarks” from subatomic physics.

Fundamentalism- A view centering around the claim that there is, or will be, one ultimate scientific theory that provides an adequate description of all phenomena everywhere.

Intertheory Relations- A comparison made between two scientific theories. Intra-level reductions and Inter-level reductions are both subclasses of this category.

Intra-level Reduction- A reduction between two theories that each provide competing descriptions of the same type of phenomena, yet on slightly different terms. These reductions occur from within the same “branch of science”, such as occurring within physics or within biology. This correlates with how Wimsatt employs the terminology (2006). *Example:* Reducing a Type I String theory to M theory.

Inter-level Reduction - A reduction between two theories that seeks to describe the same theory on different terms/methodologies. These reductions occur from different “branches of science”,

such as between chemistry and physics. This correlates with how Wimsatt employs the terminology (2006). *Example*: Reducing thermodynamics to statistical mechanics.

Pluralism – A view that rejects the fundamentalist's assertion of a final singular theory, claiming instead that many theories operating at many different are necessary to describe the world.

Physicalism- Physics is the only true nomological and ontological guide to the world. Other theories and entities are mere approximations or conglomerates; they are secondary sciences while physics is primary.

Reduction Model – An often-formalized account of the conditions involved in claiming that a reduction has taken place. *Example*: Schaffner's Model (§2.1).

Scientific Pyramid – A picture of science that organizes science into different levels, on the basis of how the entities of study are composed. Typically micro-physics and chemistry form the bottom of the pyramid, and macro-biology and psychology form the top. One of the first formulations was due to Putnam and Oppenheim (1958). See also Cartwright (1999, 7).

Successional Reduction- An intra-level reduction between an older, once-accepted theory and a newer, later accepted theory. *Example*: Reducing Mendelian inheritance to contemporary genetics.

§B Bibliography

- Adams, E. 1959. "The Foundations of Rigid Body Mechanics and the Derivation of Its Laws from Those of Particle Mechanics." In , 250–65. North-Holland.
- Airy, George Biddell. 1838. "On the Intensity of Light in the Neighborhood of a Caustic." *Transactions of the Cambridge Philosophical Society* VI: 379–403.
- Alison Wylie. 2000. "Questions of Evidence, Legitimacy, and the (Dis)Unity of Science." *American Antiquity* 65 (2): 227–37.
- Bahcall, John N., and Giora Shaviv. 1968. "Solar Models and Neutrino Fluxes." *The Astrophysical Journal* 153 (July): 113. doi:10.1086/149641.
- Balzer, W., C.U. Moulines, and J.D. Sneed. 1987. *An Architectonic for Science: The Structuralist Program*. 1st ed. Springer.
- Batterman, Robert W. 1995. "Theories between Theories: Asymptotic Limiting Intertheoretic Relations." *Synthese* 103 (2): 171–201. doi:10.2307/20117397.
- . 1997. "'Into a Mist': Asymptotic Theories on a Caustic." *Studies in History and Philosophy of Science Part B* 28 (3): 395–413.
- . 2002a. *The Devil in the Details: Asymptotic Reasoning in Explanation, Reduction, and Emergence*. Oxford University Press.
- . 2002b. *The Devil in the Details: Asymptotic Reasoning in Explanation, Reduction, and Emergence*. Oxford University Press.
- . 2005. "Response to Belot's 'Whose Devil? Which Details?.'" *Philosophy of Science* 72 (1): 154–63.
- Beard, Kenneth V., and Catherine Chuang. 1987. "A New Model for the Equilibrium Shape of Raindrops." *Journal of the Atmospheric Sciences* 44 (11): 1509–24. doi:10.1175/1520-0469(1987)044<1509:ANMFTE>2.0.CO;2.
- Belot, Gordon. 2000. "Chaos and Fundamentalism." *Philosophy of Science* 67 (September): S454–65.
- . 2005. "Whose Devil? Which Details?" *Philosophy of Science* 72 (1): 128–53. doi:10.1086/428072.

- Bender, Carl M., and Steven A. Orszag. 1978. *Advanced Mathematical Methods for Scientists and Engineers: Asymptotic Methods and Perturbation Theory*. Springer.
- Bickle, John. 1992. "Revisionary Physicalism." *Biology and Philosophy* 7 (4): 411–30.
- . 1993. "Connectionism, Eliminativism, and the Semantic View of Theories." *Erkenntnis* (1975-) 39 (3): 359–82.
- . 1998. *Psychoneural Reduction the New Wave*. Cambridge, Mass.: MIT Press.
- Bobis, Laurence, and James Lequeux. 2008. "Cassini, Rømer, and the Velocity of Light." *Journal of Astronomical History and Heritage* 11 (July): 97–105.
- Bohr, Niels. 1913. "On the Constitution of Atoms and Molecules, Part I." *Philosophical Magazine* 26 (151): 1–25.
- Bortfeldt, Jürgen. 1992. *Fundamental Constants in Physics and Chemistry*. Springer.
- Callaway, Ewen, and Nature magazine. 2012. "Carbon Dating Gets a Reset: Scientific American." October 12. <http://www.scientificamerican.com/article.cfm?id=carbon-dating-gets-reset>.
- Cartwright, Nancy. 1994. "Fundamentalism vs. the Patchwork of Laws." *Proceedings of the Aristotelian Society*, New Series, 94: 279–92.
- . 1999. *The Dappled World : A Study of the Boundaries of Science*. Cambridge University Press.
- Chang, Hasok. 2007. *Inventing Temperature: Measurement and Scientific Progress*. Oxford: Oxford University Press.
- Churchland, Patricia Smith. 1989. *Neurophilosophy: Toward a Unified Science of the Mind-Brain*. Cambridge, Mass.: MIT Press.
- Churchland, Paul M. 1985. "Reduction, Qualia, and the Direct Introspection of Brain States." *The Journal of Philosophy* 82 (1): 8–28. doi:10.2307/2026509.
- Daukantas, Patricia. 2009. "Ole Römer and the Speed of Light." *Optics and Photonics News* 20 (7): 42–47. doi:10.1364/OPN.20.7.000042.
- Descartes, René. 2001. *Discourse on Method, Optics, Geometry, and Meteorology*. Hackett Publishing.
- Dizadji-Bahmani, Foad, Roman Frigg, and Stephan Hartmann. 2010. "Who's Afraid of Nagelian Reduction?" *Erkenntnis* 73 (3): 393–412.

- Eck, Dingmar van, Huib Looren De Jong, and Maurice K. D. Schouten. 2006. "Evaluating New Wave Reductionism: The Case of Vision." *The British Journal for the Philosophy of Science* 57 (1): 167–96.
- Endicott, Ronald. 2001. "Post-Structuralist Angst-Critical Notice: John Bickle, Psychoneural Reduction: The New Wave." *Philosophy of Science* 68 (3): 377–93.
- Endicott, Ronald P. 1998. "Collapse of the New Wave." *The Journal of Philosophy* 95 (2): 53–72. doi:10.2307/2564571.
- . 2007. "Reinforcing the Three 'R'S: Reduction, Reception, and Replacement." In *The Matter of the Mind: Philosophical Essays on Psychology, Neuroscience, and Reduction*, edited by M. Schouten and H. Looren de Jong. Blackwell.
- Feyerabend, Paul K. 1962. "Explantion, Reduction and Empiricism." In *Minnesoda Studies in the Philosophy of Science*, 28–97. University of Minnesota Press.
- Fine, Arthur. 1996. *The Shaky Game*. 1st ed. University Of Chicago Press.
- Fraassen, Bas. C. van. 1980. *The Scientific Image*. Oxford University Press, USA.
- Geroch, Robert. 1981. *General Relativity from A to B*. University of Chicago Press.
- Goddard, William A. 1985. "Theoretical Chemistry Comes Alive: Full Partner with Experiment." *Science*, New Series, 227 (4689): 917–23.
- Grev, R.S., and H.F. Schaefer. 1985. *Journal of Chemistry* 82 (4126).
- Hempel, Carl. 1966. *Philosophy of Natural Science*. Prentice Hall.
- Hempel, Carl G. 1970. *Aspects of Scientific Explanation: And Other Essays in the Philosophy of Science*. 5th Printing. Free Press.
- Hentschel, Klaus. 2009. "Atomic Models, Nagaoka's Saturnian Model." In *Compendium of Quantum Physics*, edited by Daniel Greenberger, Klaus Hentschel, and Friedel Weinert, 22–23. Springer Berlin Heidelberg. http://link.springer.com/chapter/10.1007/978-3-540-70626-7_10.
- Hershey, A. D., and Martha Chase. 1952. "Independent Functions of Viral Protein and Nucleic Acid in Growth of Bacteriophage." *The Journal of General Physiology* 36 (1): 39–56. doi:10.1085/jgp.36.1.39.
- Hooker, C.a. 1981a. "Towards a General Theory of Reduction. Part I: Historical and Scientific Setting." *Dialogue: Canadian Philosophical Review/Revue Canadienne de Philosophie* 20 (01): 38–59. doi:10.1017/S0012217300023088.

- . 1981b. “Towards a General Theory of Reduction. Part II: Identity in Reduction.” *Dialogue: Canadian Philosophical Review/Revue Canadienne de Philosophie* 20 (02): 201–36. doi:10.1017/S0012217300023301.
- Hoyningen-Huene, Paul, and Franz M Wuketits. 1989. *Reductionism and Systems Theory in the Life Sciences Some Problems and Perspectives*. Dordrecht: Springer Netherlands.
- Jackson, J. D. 1999. “From Alexander of Aphrodisias to Young and Airy.” *Physics Reports* 320 (1–6): 27–36. doi:10.1016/S0370-1573(99)00088-5.
- Jaeger, Gregg. 1998. “The Ehrenfest Classification of Phase Transitions: Introduction and Evolution.” *Archive for History of Exact Sciences* 53 (1): 51–81. doi:10.1007/s004070050021.
- Jech, Thomas J. 1978. *Set Theory*. Academic Press.
- John A. Adam. 2002. “The Mathematical Physics of Rainbows and Glories.” *Physics Reports*, no. 356: 229–365.
- Kafafi, Z.H., R.H. Hague, L. Fredin, and J.L. Margrave. 1983. *Journal of Chemical Physics* 87 (787).
- Kellert, Stephen H, Helen E Longino, and C. Kenneth Waters. 2006. *Scientific Pluralism*. Minneapolis, MN: University of Minnesota Press.
- Kemeny, John G., and Paul Oppenheim. 1956. “On Reduction.” *Philosophical Studies: An International Journal for Philosophy in the Analytic Tradition* 7 (1/2): 6–19.
- Kitcher, Philip, and Wesley C Salmon. 2011. *Scientific Explanation*. [Minneapolis]: University of Minnesota Press.
- Kuhn, Thomas S. 1962. *The Structure of Scientific Revolutions*. University of Chicago Press.
- Laudan, Larry. 1981. “A Confutation of Convergent Realism.” *Philosophy of Science* 48 (1): 19–49.
- Lederman, Leon M., and Christopher T. Hill. n.d. *Quantum Physics for Poets*. Prometheus Books.
- Malament, David B. 1986a. “Newtonian Gravity, Limits, and the Geometry of Space.” In *From Quarks to Quasars: Philosophical Problems of Modern Physics*, 181–201. University of Pittsburg Series in the Philosophy of Science 7. University of Pittsburgh Press.

- . 1986b. “Gravity and Spatial Geometry.” *Logic, Methodology and Philosophy of Science, Proceedings of the Seventh International Congress of Logic, Methodology, and Philosophy of Science, Salzburg, 1983*, no. 117: 405–11.
- . 1995. “Is Newtonian Cosmology Really Inconsistent?” *Philosophy of Science* 62 (4): 489–510.
- . 2012. *Topics in the Foundations of General Relativity and Newtonian Gravitational Theory*. University of Chicago Press.
<http://www.lps.uci.edu/malament/FndsofGR/GR.pdf>.
- Michelson, Albert Abraham, and Edward Morley. 1887. “On the Relative Motion of the Earth and the Luminiferous Ether.” *American Journal of Science* 34 (203): 333–45.
- Mormann, Thomas. 1988. “Structuralist Reduction Concepts as Structure-Preserving Maps.” *Synthese* 77 (2): 215–50.
- Nagel, Ernest. 1935. “The Logic of Reduction in the Sciences.” *Erkenntnis* 5 (January): 46–52.
- . 1951. “Mechanistic Explanation and Organismic Biology.” *Philosophy and Phenomenological Research* 11 (3): 327–38. doi:10.2307/2103537.
- . 1979. *The Structure of Science: Problems in the Logic of Scientific Explanation*. 2nd ed. Hackett Publishing Company.
- Newton, Sir Isaac. 1718. *Opticks: Or, A Treatise of the Reflections, Refractions, Inflections and Colours of Light*. ... Printed for W. and J. Innys, printers to the Royal Society, at the Prince’s-Arms in St. Paul’s Church-Yard.
- Nickles, Thomas. 1973. “Two Concepts of Intertheoretic Reduction.” *The Journal of Philosophy* 70 (7): 181–201.
- Niebergall, Karl-Georg. 2002. “Structuralism, Model Theory and Reduction.” *Synthese* 130 (1): 135–62.
- Oppenheim, Paul, and Hilary Putnam. 1958. “Unity of Science as a Working Hypothesis.” In , edited by Herbert Feigl, Michael Scriven, and Grover Maxwell, 2:University of Minnesota Press – 3.
- Presilla-Marquez, J.D., and W.R. Graham. 1991. *Journal of Chemical Physics* 95 (5612).
- Pruppacher, H. R., and J. D. Klett. 2010. *Microphysics of Clouds and Precipitation*. Springer.
- Putnam, Hilary. 1965. “How Not to Talk About Meaning.” In *Boston Studies in the Philosophy of Science, Vol. II*,. Vol. 2. Humanities Press.

- Ramsey, Christopher Bronk, Richard A. Staff, Charlotte L. Bryant, Fiona Brock, Hiroyuki Kitagawa, Johannes van der Plicht, Gordon Schlolaut, et al. 2012. "A Complete Terrestrial Radiocarbon Record for 11.2 to 52.8 Kyr B.P." *Science* 338 (6105): 370–74. doi:10.1126/science.1226660.
- Redhead, M. 2004. "Asymptotic Reasoning." *Studies in History and Philosophy of Science Part B* 35 (3): 527–530.
- Richardson, Robert C. 2008. "Autonomy and Multiple Realization." *Philosophy of Science* 75 (5): 526–36. doi:10.1086/598956.
- Riel, Raphael van. 2011. "Nagelian Reduction beyond the Nagel Model." *Philosophy of Science* 78 (3): 353–75.
- Rivadulla, Andrés. 2004. "The Newtonian Limit of Relativity Theory and the Rationality of Theory Change." *Synthese* 141 (3): 417–29.
- Rohrlich, Fritz. 1988. "Pluralistic Ontology and Theory Reduction in the Physical Sciences." *The British Journal for the Philosophy of Science* 39 (3): 295–312.
- Rohrlich, Fritz, and Larry Hardin. 1983. "Established Theories." *Philosophy of Science* 50 (4): 603–17. doi:10.2307/187558.
- Rovelli, Carlo. 2013. "Aristotle's Physics: A Physicist's Look." *arXiv:1312.4057 [physics]*, December. <http://arxiv.org/abs/1312.4057>.
- . n.d. "Aristotle's Physics."
- Salmon, Wesley C. 2006. *Four Decades of Scientific Explanation*. University of Pittsburgh Pre.
- Sarkar, Sahotra. 1992. "Models of Reduction and Categories of Reductionism." *Synthese* 91 (3): 167–94.
- Scerri, Eric R. 1994. "Has Chemistry Been at Least Approximately Reduced to Quantum Mechanics?" *PSA: Proceedings of the Biennial Meeting of the Philosophy of Science Association* 1994: 160–70.
- Schaffner, Kenneth F. 1967. "Approaches to Reduction." *Philosophy of Science* 34 (2): 137–47.
- . 1974. "Reductionism in Biology: Prospects and Problems." *PSA: Proceedings of the Biennial Meeting of the Philosophy of Science Association* 1974: 613–32.
- . 1977. "Reduction, Reductionism, Values and Progress in the Biomedical Sciences." In *The Pittsburgh Series in the Philosophy of Science*, 143–71.

- . 2006. “Reduction: The Cheshire Cat Problem and a Return to Roots.” *Synthese* 151 (3): 377–402.
- Sklar, Lawrence. 1967. “Types of Inter-Theoretic Reduction.” *The British Journal for the Philosophy of Science* 18 (2): 109–24.
- . 1995. *Physics and Chance: Philosophical Issues in the Foundations of Statistical Mechanics*. Cambridge University Press.
- Stegmüller, Wolfgang. 1976. *The Structure and Dynamics of Theories*. Springer-Verlag.
- Suarez, Mauricio, ed. 2008. *Fictions in Science: Philosophical Essays on Modeling and Idealization*. 1st ed. Routledge.
- Suppes, Patrick. 1957. *Introduction to Logic*. Vol. 10. 40. Dover Publications.
- . 1967. “What Is a Scientific Theory?” In *Philosophy of Science Today*, 55–67. Basic Books Inc.
- Thomson, J.J. 1904. “XXIV. On the Structure of the Atom: An Investigation of the Stability and Periods of Oscillation of a Number of Corpuscles Arranged at Equal Intervals around the Circumference of a Circle; with Application of the Results to the Theory of Atomic Structure.” *Philosophical Magazine Series 6* 7 (39): 237–65.
doi:10.1080/14786440409463107.
- Trautman, Andrzej. 1965. “Foundations and Current Problem of General Relativity.” In *Space-Time*. Englewood Cliffs, NJ: Prentice Hall.
- Weatherall, James Owen. 2011. “On (Some) Explanations in Physics.” *Philosophy of Science* 78 (3): 421–47.
- Weinberg, Steven. 1994. *Dreams of a Final Theory: The Scientist’s Search for the Ultimate Laws of Nature*. Vintage.
- Weisstein, Eric W. 2015a. “Totally Ordered Set.” Text. Accessed January 20.
<http://mathworld.wolfram.com/TotallyOrderedSet.html>.
- . 2015b. “Partial Order.” Text. Accessed January 20.
<http://mathworld.wolfram.com/PartialOrder.html>.
- Wimsatt, William C. 1974. “Reductive Explanation: A Functional Account.” *PSA: Proceedings of the Biennial Meeting of the Philosophy of Science Association* 1974: 671–710.
- . 2006. “Reductionism and Its Heuristics: Making Methodological Reductionism Honest.” *Synthese* 151 (3): 445–75.

- . 2007. *Re-Engineering Philosophy for Limited Beings: Piecewise Approximations to Reality*. annotated edition. Harvard University Press.
- Winsberg, Eric. 2006. “Handshaking Your Way to the Top: Simulation at the Nanoscale.” *Philosophy of Science* 73 (5): 582–94. doi:10.1086/523778.
- Winther, Rasmus Grønfeldt. 2009. “Schaffner’s Model of Theory Reduction: Critique and Reconstruction.” *Philosophy of Science* 76 (2): 119–42. doi:10.1086/603620.
- Woodger, J. H. 1952. *Biology & Language*. Cambridge University Press.
- Young, Thomas. 1804. “The Bakerian Lecture: Experiments and Calculations Relative to Physical Optics.” *Philosophical Transactions of the Royal Society of London* 94 (January): 1–16.