

Zero or minimum degree? Rethinking minimum gradable adjectives¹

Ciyang QING — *Stanford University*

Abstract. This paper examines the class of minimum gradable adjectives (minimum GAs), the standards of whose positive forms are traditionally characterized in terms of minimum degrees. Using *profitable* as a case study, I argue that their standards are better characterized using zero degrees, which may not be minimum degrees. Moreover, I show that this zero-standard interpretation cannot be derived by existing approaches, and propose that it is derived by a new compositional mechanism, which also derives the truth conditions of plain comparatives and is different from the one that derives the relative-standard interpretation of a positive form.

Keywords: Gradable adjectives, minimum-standard adjectives, comparatives, zero degrees

1. Introduction

This paper examines the class of minimum(-standard) gradable adjectives (e.g., *wet* and *bent*), whose positive forms have an interpretation that is traditionally linked to the minimum degree on the adjective's scale (e.g., Kennedy and McNally, 2005; Kennedy, 2007). For instance, the scale of *wet* has a minimum degree, which corresponds to zero amount of wetness (i.e., completely dry). Correspondingly, *wet* has a minimum standard: something is wet iff it has non-zero amount of wetness.² Within a degree-semantic framework, the interpretation of the positive form of a minimum adjective in context can be formally represented in (1). Here I intentionally represent only the surface form, so that I am not presupposing any specific analysis that assumes silent materials.

$$(1) \quad \llbracket x \text{ is } A_{\min} \rrbracket^c = \mathbf{A}\text{-measure}(x) > d_{\min}$$

Note that the formal representation of the interpretation of the positive form of a minimum adjective makes reference to the minimum degree (1). Meanwhile, informally we also say that something is wet iff it has non-zero amount of wetness, which strictly speaking characterizes the standard of *wet* in terms of the zero degree. In the case of *wet*, these two descriptions are equivalent because the minimum degree is also a meaningful, non-arbitrary zero point on the scale, i.e., it indicates the absence of the relevant property. In fact, this is the case for so many of the minimum adjectives, that zero and minimum degrees are often used interchangeably in the literature. However, in this paper, I argue that zero and minimum degrees are not always identical, and that the standard of the so-called “minimum interpretation” is in fact based on the zero degree rather than the minimum. In other words, the class of *minimum-standard gradable adjectives* in fact have standards that are zero degrees rather than minimum degrees.

Whether the standard of a minimum GA should be characterized as a zero or minimum degree also has broader theoretical implications concerning how to distinguish between analyses of GAs that postulate semantic underspecification and those that postulate ambiguity. In general,

¹I would like to thank Dan Lassiter, Cleo Condoravdi, and Judith Degen for the supervision that greatly helped improve this work. I would also like to thank Chris Barker, Louis McNally, Stephanie Solt for the very helpful discussions, and SuB reviewers for their insightful comments and suggestions. All errors are my own.

²Note that we need to factor out granularity/imprecision, e.g., a towel only has a few drops of water on it may not be considered wet, either because people cannot tell the difference from a completely dry towel, or because they do not care about it as far as the contextual goal is concerned.

if we observe variability in the interpretations of an expression, or in this case a class of expressions, a theory that posits semantic underspecification needs to be supplemented with a contextual-resolution mechanism that accounts for the observed variability, and a theory that posits ambiguity should ideally be supported by independent evidence. In this paper, I illustrate that for GAs that do not have minimum degrees on their scales, their zero-standard interpretations are in fact not derivable from major existing analyses of GAs that postulate semantic underspecification. Instead, I propose an ambiguity analysis of minimum GAs, according to which the zero-standard and relative-standard interpretations of their positive forms are derived by two independent compositional mechanisms. I further argue that this ambiguity analysis provides a better account of the similarities and differences between positive forms and comparative constructions.

The rest of the paper is organized as follows. In the next section, I will use the gradable adjective *profitable* as a case study, and argue that it has a meaningful zero degree as its standard, which is not a minimum degree on its scale. This suggests that it is descriptively more accurate to call such an interpretation a > 0 interpretation. Moreover, I argue that this interpretation cannot be derived from previous analyses that assume semantic underspecification. This provides some initial evidence in favor of an ambiguity analysis, which is further motivated in section 3 by considering the similarities and differences between positive forms and comparative constructions. In section 4, I provide a compositional semantic analysis of the > 0 interpretation, based on Schwarzschild and Wilkinson's (2002) analysis of comparative constructions. I discuss more cases that support this analysis and some remaining issues in section 5.

2. A case study: *profitable*

2.1. The > 0 interpretation of *profitable*

As a case study, let us consider the adjective *profitable*. This adjective is clearly gradable, as can be seen from the examples in (2).

- (2) a. This company is more profitable than that company.
- b. How profitable was this company last year?
- c. This industry is so profitable that everybody wants to get a share.

The positive form of *profitable* has an interpretation whose standard is 0, i.e., *x was profitable* has an interpretation which is true iff *x* at least made a minimum/non-zero amount of profit. Consider the following naturally occurring example (3).

- (3) [Headline] Spotify, the leading music streaming app, is finally profitable.
 [Main text] [...] Today, for the very first time, the company is reporting that it's turned a profit. That's right: some 13 years and 96 million paid subscribers later, Spotify is finally making money.

In the headline, the author asserted that Spotify is profitable, and from the main text it is clear that the intended interpretation is that the company has turned a profit/is making money.

This shows that *profitable* has a zero-standard interpretation (henceforth > 0 interpretation). However, depending on the context, this may not be the most salient interpretation. For concreteness, throughout this paper I will assume that the performances of six companies in 2019 are listed in (4), and that only companies *D*, *E*, and *F* are oil companies.

(4)	Company	A	B	C	D	E	F
	Profit/Loss	-\$100M	-\$10M	0	\$5K	\$1M	\$1B

The polar question (5a) can be interpreted either as asking whether Company D made any profit at all, or as asking whether Company D made an amount of profit that “stands out” among some contextual comparison class, e.g., oil companies. As a result, one can felicitously answer it with “yes” or “no,” or (5b) to avoid miscommunication.

- (5) a. Was Company D profitable in 2019?
b. Company D was (technically) $\text{profitable}_{>0}$, but not $\text{profitable}_{\geq\theta}$ for an oil company.

This example shows that *profitable* can have a zero standard or a relative standard. Note that unlike the relative interpretation of *profitable*, the > 0 interpretation is not vague. In this respect, this > 0 interpretation of *profitable* patterns with run-of-the-mill minimum adjectives such as *wet* and *bent*. Moreover, just as we can use *slightly* to modify *wet* or *bent*, we can do so to modify *profitable*, e.g., in the naturally occurring example below (6).

- (6) “If they hit that number, it’s going to equate to 48,000 model 3s produced in the September quarter. That should get them to profitability, slightly profitable,” Munster said.

Following Solt’s (2012) analysis of *slightly*-modification, I take this as further evidence that the > 0 interpretation of *profitable* has a non-vague standard, i.e., the zero degree. Note that unlike Kennedy (2007), Solt (2012) does not take the compatibility with *slightly* as indicating the existence of a minimum degree, and therefore we need to independently determine the scale structure for *profitable*, by considering how the profitability of a company should be measured.

2.2. An intuitive scale structure of *profitable* and problems for previous analyses of GAs

In order to compositionally derive interpretations of various degree constructions that contain the gradable adjective *profitable*, we need to first specify the measure function it denotes (and in doing so, we would also have specified its scale structure). Intuitively, the profitability of a company can be measured by the amount of the profit it makes.³ However, there is a complication (besides the complicated actual accounting process to determine profit). It is possible for companies to lose money. The critical question is what should be the output of this measure function for companies that lost money (i.e., A and B in our working example).

I propose that we take (4) at face value, i.e., we simply treat losses as negative degrees of profitability. The resulting scale structure is shown schematically in (7). I will call this the *full-range* scale. Crucially, this is an open scale without a minimum (or maximum) degree.



This full-range scale is highly intuitive. Later I will compare it with two alternatives proposed in the literature (not specifically for *profitable*, but for similar adjectives) and argue against

³Of course, this is not the only way to measure profitability. For example, a commonly used measure is *profit margin*, which is profit divided by revenue. If we use this measure of profitability, then it is possible for a company to be more profitable than another even if the two companies made the same amount of profit. I choose profit as the measure for *profitable* just for simplicity. Nothing I say in this paper crucially hinges on this choice.

those alternatives. For now let us assume that (7) is indeed the scale structure of *profitable* and discuss its consequences. Specifically, I will consider two major types of analyses of positive forms of GAs in the literature that assume semantic underspecification and show that they both fail to derive the > 0 interpretation. These analyses all assume semantic underspecification in that they assume a single compositional mechanism that determines the standard of a positive form (8), which turns out to be different for different subclasses of GAs.

$$(8) \quad \llbracket pos \text{ Adj} \rrbracket = \lambda x. \mu_{Adj}(x) \geq \theta$$

First, I consider recent distribution-based probabilistic models (Lassiter and Goodman, 2013, 2015; Qing and Franke, 2014). Due to space constraints, I will only summarize their common assumptions and major predictions, without discussing all the details of these models and the differences between them. According to such models, the standard of a GA is calculated based on a probability distribution k over degrees, whose shape reflects our world knowledge of how the relevant degrees are distributed. For instance, in a context where we are comparing John with the general US adult male population, the standard of *John is tall* is calculated based on the probability distribution of such people's heights. Crucially, the probabilistic models make two general predictions for open-scale GAs: (i) the standard is higher than the average of the probability distribution, which accounts for the inference from *John is tall* to *John is taller than average*, and (ii) the standard is vague.

These two predictions turn out problematic if we were to use probabilistic models to derive the > 0 interpretation of *profitable*. Because of (i), we need to assume that 0 is higher than the average of probability distribution over profits. In other words, we need to assume that companies on average lost money. However, while it seems inconsistent for someone to use $6'$ as the standard of *tall* while believing that the average height is greater than $6'$, it seems perfectly consistent for someone to use 0 as the standard of *profitable* and believe that companies on average made a profit. This contrast makes it not very plausible to stipulate a probability distribution whose average is less than 0 in order to derive the > 0 interpretation of *profitable*. Even worse, even if one could motivate this assumption, because of (ii), the probabilistic models would still wrongly predict that the zero standard of *profitable* is vague. Given that (ii) is a major feature of probabilistic models, I conclude that they fail to derive the > 0 interpretation of *profitable*.

Second, I consider Kennedy's (2007) classical analysis based on Interpretive Economy (IE). This principle requires that one "maximize the contribution of the conventional meanings of the elements of a sentence to the computation of its truth conditions" (p. 36). The intended effect of IE is to ensure that if the scale of a GA has a minimum or maximum degree, then the positive form must have a minimum or maximum standard. Here I will stick with Kennedy's original proposal, which takes IE as a strict rule and makes an auxiliary assumption that maximum and minimum degrees are *natural transitions* to derive this intended effect from IE. Later I will also discuss variants of IE-based analyses and their limitations.

Crucially, according to Kennedy (2007), open scales do not have natural transitions: "More generally, there is nothing inherent to the structure of an open scale that results in natural transitions: open scales represent infinitely increasing or decreasing measures" (p. 35). Therefore, under the assumption that *profitable* has an open scale, Kennedy's original proposal would predict that it should pattern with other open-scale adjectives such as *tall* and *big* and only has a vague, relative interpretation. In other words, it fails to derive the non-vague > 0 interpretation.

Summing up the discussion so far, if the full-range scale structure of *profitable* (7) is correct, i.e., *profitable* has an open scale, then neither the probabilistic models nor the IE-based analysis can derive its non-vague > 0 interpretation. Of course, proponents of such approaches may try to avoid this problem by postulating alternative scale structures that do have minimum degrees. Below I will consider two such alternatives and argue against them.

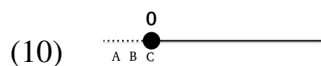
2.3. Alternative scale structures for *profitable* and arguments against them

The two alternative scale structures I will consider below are proposed in the literature not specifically for *profitable* but for similar adjectives. I will show the major problem of each of these two analyses and conclude that *profitable* is best analyzed as having the full-range scale.

The sentences I will use involve various comparative constructions, including bare comparatives (9a) or modified ones (9b). The modifier can be a measure phrase such as \$5K, or an adverb such as *slightly* or *much*. It is represented as DIFF because intuitively its meaning concerns the difference between the profits of x and y . Note that for now I abstract away from the compositional path that derives the truth conditions of (9a) and (9b), only assuming that the truth conditions are calculated based on **profit**(x), **profit**(y), and the meaning of DIFF.

- (9) a. $\llbracket x \text{ is more profitable than } y \rrbracket = \mathbf{profit}(x) > \mathbf{profit}(y)$
 b. $\llbracket x \text{ is DIFF more profitable than } y \rrbracket = \llbracket \text{DIFF} \rrbracket (\mathbf{profit}(x) - \mathbf{profit}(y))$

I will call the first analysis, proposed by Sawada and Grano (2011), the *truncation* analysis. When applied to *profitable*, it stipulates that the measure function **profit** is undefined for companies that lost money. The resulting truncated scale structure is shown schematically in (10).



A major motivation for the truncation analysis is the infelicity of (11). According to the truncation analysis, *John was later than Bill* is defined only when neither of them was early. The *but*-clause contradicts this definedness condition and therefore is infelicitous.

- (11) John was later than Bill, #but they were both early.

It is not entirely clear whether the truncation analysis provides a satisfying account even for *early* and *late*. As Sawada and Grano (2011) acknowledge, the infelicity of (11) is weak and some speakers may think it is fine. Therefore, even though I use their # annotation for this example, I will use ? to better indicate weak infelicity for future examples. In addition, I note that the infelicity seems even weaker when we use the adverb *late* instead. Multiple native speakers I have consulted reported that (12) sounded better than (11), and some found (12) perfectly fine.

- (12) John arrived later than Bill, ?but they were both early.

Moreover, there are naturally occurring examples similar to (12).

- (13) a. There was traffic around the hotel and an accident on the freeway. I got there a little later than normal, not late. My workday starts at 4:30 p.m., in my opinion. I got to the stadium at 4:04 p.m.

- b. Finance Manager Marlene Kelleher said that due to the implementation of the District's new finance software, the audit was being presented slightly later than usual but still on time.

For example, from (13a) we know that the speaker got to the stadium early, and he normally arrived even earlier. Similarly, from (13b) we can infer that the audit was usually presented early. Given that the measure function of *late* is undefined for any time point before “on time” according to the truncation analysis, it will have trouble accounting for (13). Sawada and Grano (2011: fn. 4) tentatively suggest that speakers who accept (11) are resetting the minimum value of the scale so that the scale includes enough time points before “on time,” and presumably they would say the same for (13).

Setting aside whether this additional stipulation really provides a satisfying explanation, let us focus on consequences of applying the truncation analysis to *profitable*. Since Company A lost money, the measure function of *profitable* is undefined for it, and therefore all sentences in (14) are predicted to be infelicitous due to undefinedness (their truth conditions are shown on the right for easy reference).

- | | | | |
|------|----|--|---|
| (14) | a. | ? Company B was more profitable than A. | $\mathbf{profit(b)} > \mathbf{profit(a)}$ |
| | b. | Company E was more profitable than A. | $\mathbf{profit(e)} > \mathbf{profit(a)}$ |
| | c. | Company E was much more profitable than A. | $\mathbf{profit(e)} - \mathbf{profit(a)} \geq \theta_{\text{much}}$ |
| | d. | Company E was \$101M more profitable than A. | $\mathbf{profit(e)} - \mathbf{profit(a)} \geq \101M |

However, while there might be something slightly weird about (14a), (14b–14d) are clearly true.⁴ Therefore, the truncation analysis wrongly predicts (14b–14d) to be infelicitous. Note that the option of resetting the minimum value of the scale will not explain the contrast between (14a) and (14b). In order for (14b–14d) to be felicitous, we need to be able to reset the minimum value of the scale low enough so that the measure function is defined for Company A. However, if we do that, then the measure function will also be defined for Company B, which means that (14a) is now defined and we lose the initial explanation of why it is infelicitous.

In comparison, the full-range analysis correctly predicts that (14b–14d) are true. It also predicts that (14a) is true. While this leaves the infelicity or weirdness of (14a) unexplained, we can reasonably assume that this is in fact the correct prediction as far as the semantics is concerned. Similar to the discussion above with *early/late* and as indicated by the ? annotation, the infelicity of (14a) is weak, and the dominant intuition reported by the native speakers I have consulted is that (14a) is technically true but very misleading. This strongly suggests that the infelicity of (14a) should be explained in the pragmatics rather than in the semantics. Therefore, I conclude that the full-range analysis, when supplemented with a proper pragmatic theory, should be preferred to the truncation analysis.

Now I turn to the second alternative, proposed by Bierwisch (1989), and call it the *compression* analysis. When applied to *profitable*, it stipulates that the measure function **profit** maps all the

⁴For those who need more convincing that (14b) is true and felicitous. Here is a naturally occurring example: “Did you know that your business is more profitable than Amazon.com? Literally. [...] But really and truly, your business is making more money than Amazon.com. In fact, just last quarter, the company lost \$126 million and expects to have an operating loss of \$810 million this year. [...] So congratulations: your business is profitable and has been for years. Amazon’s is not.”

companies that lost money (as well as those that broke even) to the zero degree. The resulting compressed scale structure is shown schematically in (15).



Bierwisch (1989) proposes the scale structure in (15) to analyze what he calls *evaluative* adjectives such as *pretty/ugly* and *industrious/lazy*. According to him, an evaluative adjective is not directly associated with a measure function. For instance, *pretty* simply denotes an individual property *P* (which may be context-sensitive), and it becomes gradable only relative to a comparison class. The individuals in the comparison class are ordered wrt the degree or extent to which they satisfy *P*. In particular, all the individuals that do not satisfy *P* are mapped (or compressed) to the same zero degree.

In the case of *profitable*, setting aside whether its gradability really relies on comparison classes, the compression analysis correctly predicts that (14b) is true. However, it predicts that (14a) is false. Given the discussion above, this does not seem to be the right way to capture its (weak) infelicity, but it is not obviously wrong, either. Its predictions for (14c), however, is much worse. To highlight the problem, note that it predicts that the sentence pairs in (16) and (17) have the same truth values, because the measure function maps A and C to the same value 0. However, these predictions are clearly wrong. We can reasonably say that (16a) is true and (16b) is false, and similarly that (17a) is false and (17b) is true. Finally, (14d) should be true, but the compression analysis wrongly predicts that it is false.

- (16) a. D was (only) slightly more profitable than C. $0 < \mathbf{profit(d)} - \mathbf{profit(c)} \leq \theta_{\text{slightly}}$
b. D was (only) slightly more profitable than A. $0 < \mathbf{profit(d)} - \mathbf{profit(a)} \leq \theta_{\text{slightly}}$
- (17) a. Company D was much more profitable than C. $\mathbf{profit(d)} - \mathbf{profit(c)} \geq \theta_{\text{much}}$
b. Company D was much more profitable than A. $\mathbf{profit(d)} - \mathbf{profit(a)} \geq \theta_{\text{much}}$

Therefore, we have seen that the compression analysis as stated above makes incorrect predictions when comparatives are modified by adverbs or measure phrases.

I should note that this compression analysis is actually not the full theory proposed by Bierwisch (1989). According to him, while an evaluative adjective primarily has the compressed scale structure in (15), some evaluative adjectives can also have the full-range scale structure in (7) by conjoining the two compressed scales of the evaluative adjective and its antonym “back to back” at the zero point. In effect, Bierwisch is proposing a hybrid account that assumes both the full-range scale and the compressed scale.

Nevertheless, in the case of *profitable*, I argue that the full-range analysis (supplemented with a proper pragmatic analysis) should still be preferred to such a hybrid account. First, since the hybrid account also needs to be supplemented with a pragmatic analysis to spell out which scale structure to use in any given context, all else being equal, the full-range analysis should be preferred because it is semantically more parsimonious. Second, note that the full-range scale structure in (7) is needed in order to derive the correct truth values for (14d), (16a), and (17b), and that the compressed scale structure generates wrong truth values for such cases. It would be very difficult to explain why the full-range scale structure is obligatorily used for such

cases, in particular if the compressed scale structure is assumed to be the primitive.⁵ Therefore, I conclude that the scale structure of *profitable* is best analyzed as the full-range scale in (7).

To sum up, we have seen that *profitable* arguably has a full-range scale structure that does not have a minimum degree. Nevertheless, it has a non-vague > 0 interpretation, which uses the zero degree as the standard and cannot be derived from distribution-based probabilistic models or the IE-based analysis in Kennedy's (2007) original formulation.

However, as readers might have noticed, the reason why Kennedy's (2007) original IE-based analysis fails to derive the > 0 interpretation of *profitable* has little to do with IE per se, but rather the auxiliary assumption that open scales never have natural transitions. Perhaps this assumption is too strong and can be weakened so that (at least some) zero degrees on open scales are also natural transitions. I have three responses to this attempt to salvage an IE-based analysis. First, this essentially agrees that the class of minimum-standard GAs should be better characterized as zero-standard GAs. Second, given that *profitable* also has relative interpretations, this attempt would also need to weaken IE so that it is just a preference rather than a strict rule, i.e., relative interpretations are in principle available but generally dispreferred. For proponents of IE, this move is presumably needed anyways, since there are many other independent counterexamples showing that the original formulation is too strong (e.g., Lassiter, 2017). Finally, as I will discuss in the next section, even though this weaker version of IE-based analysis (henceforth weak-IE) can account for the > 0 and relative interpretations of positive forms of GAs, it faces further challenges when we also take into account comparatives. I will argue there that it would be better to treat the > 0 interpretation as an independent reading and provide a separate compositional derivation.

3. Connections to comparative constructions

In English, measure phrases such as *3 inches* can appear with a plain gradable adjective, as well as in comparative constructions (18).⁶

- (18) a. This glass is 3 inches tall. (plain)
 b. This glass is 3 inches taller than that glass. (comparative)

Whether a measure phrase can appear with a plain gradable adjective is generally idiosyncratic both within a language and crosslinguistically. For instance, Schwarzschild (2005) observes that in English we cannot say *5 pounds heavy* (cf. 18a), but we can say so in other languages such as Italian and German. In contrast, measure phrases are generally more compatible with comparative constructions. Crucially, as Schwarzschild (2005) observes, there seems to be a universal that if a language allows measure phrases to appear with a plain adjective (he calls such measure phrases *direct* measure phrases), then it also allows them in the corresponding comparative constructions (he calls such measure phrases *indirect* measure phrases).

Note that the reverse of this universal generally does not hold. For instance, in English we can say *this bag is 5 pounds heavier* but not *this bag is 5 pounds heavy*. However, Sawada

⁵This is particularly so for (16b), which is predicted to be true based on the compressed scale structure, but is nevertheless false. This means that we cannot say that the full-range scale is used to salvage an otherwise false/infelicitous sentence.

⁶I use the term "plain" instead of "positive form of" to avoid the implication that a silent morpheme *pos* is present.

and Grano (2011) observes that this reversal does seem to hold for the class of minimum adjectives. For instance, even though direct measure phrases are uniformly banned for plain lower-open-scale adjectives in Japanese, Spanish, Korean and Russian, they are compatible with comparatives, as well as plain minimum adjectives such as *bent*.⁷

These two universals suggest that there is something about minimum adjectives that is shared by comparative constructions, which lower-open-scale adjectives do not share. This provides some initial motivation to look into comparative constructions more closely to see whether there is something that can inspire our analysis of minimum adjectives.

Sawada and Grano (2011), building on Svenonius and Kennedy (2006), propose that minimum adjectives and comparatives both have scales with a minimum degree. Specifically, they propose that comparative constructions have a scale structure with a minimum degree because the *than*-phrase transforms the original scale of the gradable adjective to one that starts with the degree introduced by the *than*-phrase (19).

- (19) $\llbracket \text{taller than Mary} \rrbracket = \lambda x. \mathbf{height}_{\mathbf{height(m)}}^{\uparrow}(x)$, where $\mathbf{height}_{\mathbf{height(m)}}^{\uparrow}$ is a transformed scale of height whose starting point is Mary's height.

The semantic type of (19) is $e \rightarrow d$. Since Sawada and Grano (2011) assume that the type of a measure phrase such as *3 inches* is d , (19) cannot be directly composed with the measure phrase. To resolve this type mismatch, Sawada and Grano (2011) assume that a null degree morpheme *Meas*, defined in (20a), first transforms (19) to type $d \rightarrow et$ (20b). This is the correct type to allow for the composition with the measure phrase, which yields the individual property in (20c).

- (20) a. $\llbracket \text{Meas} \rrbracket = \lambda g_{ed} : \underline{g \text{ has a minimum degree on the scale. } \lambda d \lambda x. g(x) \geq d}$
 b. $\llbracket \text{Meas taller than Mary} \rrbracket = \lambda d \lambda x. \mathbf{height}_{\mathbf{height(m)}}^{\uparrow}(x) \geq d$
 c. $\llbracket 3 \text{ inches Meas taller than Mary} \rrbracket = \lambda x. \mathbf{height}_{\mathbf{height(m)}}^{\uparrow}(x) \geq 3'$

Sawada and Grano (2011) further assume that *Meas* has a selectional restriction so that it is only compatible with measure functions that have a minimum degree on the scale, i.e., the underlined part in (20a). Comparatives and minimum adjectives satisfy this requirement, and therefore they are compatible with measure phrases. Lower-open-scale adjectives do not satisfy this requirement, and therefore they are generally not compatible with measure phrases. This accounts for the fact that in some languages measure phrases are totally incompatible with open-scale adjectives. However, other languages, such as English, may allow for some idiosyncratic set of open-scale adjectives to override this restriction. In this way, Sawada and Grano (2011) can account for the crosslinguistic generalizations regarding the compatibility between measure phrases and different types of degree constructions.

⁷There is a further point of variation that Sawada and Grano (2011) discuss. Measure phrases can in fact appear with plain adjectives in Japanese, but in such cases they must be interpreted as comparatives with implicit standards (unlike the other three languages, where even such interpretations are impossible). In fact, similar phenomena can be observed in English. For example, while it is generally weird to say *John is 5'5" short*, we can say *the sleeves are 1 inch short*. However, it does not mean that the length of the sleeves is 1 inch, but rather that the sleeves are 1 inch too short. I will follow Sawada and Grano (2011) in assuming that in such cases the plain adjectives are coerced to comparative interpretations.

The remaining issue Sawada and Grano need to address is how to derive the truth conditions for plain comparatives, e.g., *John is taller than Mary*. They propose that the derivation is totally in parallel with the positive form construction *John is tall*, where a silent *pos*, as defined by Kennedy (2007), is used to introduce an appropriate threshold (21a).

- (21) a. $\llbracket pos \rrbracket^C = \lambda g. \lambda x. g(x) \geq \mathbf{s}(g)(C)$, where $\mathbf{s}(g)(C)$ either “stands out” wrt g (i.e., is a natural transition) or “stands out” in C
 b. $\llbracket pos \text{ tall} \rrbracket^C = \lambda x. \mathbf{height}(x) \geq \theta_C$
 c. $\llbracket pos \text{ taller than Mary} \rrbracket^C = \lambda x. \mathbf{height}_{\mathbf{height}(\mathbf{m})}^\uparrow(x) > 0$,
 which returns true if $\mathbf{height}(x) > \mathbf{height}(\mathbf{m})$

Since *tall* has an open scale and presumably does not have a natural transition, it has a relative standard that “stands out” in the context (21b). In contrast, since *taller than Mary* has a minimum degree, i.e., 0, on its scale, which is assumed to be a natural transition, it is used as the standard (21c), and therefore we seem to have derived the correct truth conditions for plain comparatives.⁸ However, since Sawada and Grano intend $\mathbf{height}_{\mathbf{height}(\mathbf{m})}^\uparrow$ to be a truncation of the height scale, this function is undefined for people shorter than Mary (recall the earlier discussion about their analysis of *early/late*). But if John is shorter than Mary, then *John is (3 inches) taller than Mary* should be false rather than a presupposition failure. Therefore, (20c) and (21c) are in fact incorrect given how Sawada and Grano define $\mathbf{height}_{\mathbf{height}(\mathbf{m})}^\uparrow$.

Note that Sawada and Grano need to define $\mathbf{height}_{\mathbf{height}(\mathbf{m})}^\uparrow$ as a truncated scale only because they adopt Kennedy’s (2007) original proposal and have to ensure that $\mathbf{height}_{\mathbf{height}(\mathbf{m})}^\uparrow$ has a minimum degree. However, once we adopt weak-IE, we no longer need to make sure that the modified scale has a minimum degree and can simply define it as measuring the degree difference between the subject and the standard (22a). This scale has the zero degree as the natural transition, and therefore we can derive the truth conditions of plain comparatives as > 0 interpretations (22b). In particular, (22b) correctly predicts that *John is taller than Mary* is false when John is in fact shorter than Mary.

- (22) a. $\llbracket \text{taller than Mary} \rrbracket = \lambda x. \mathbf{height}(x) - \mathbf{height}(\mathbf{m})$ (cf. Sassoon, 2010)
 b. $\llbracket pos \text{ taller than Mary} \rrbracket^C = \lambda x. \mathbf{height}(x) - \mathbf{height}(\mathbf{m}) > 0$,
 which is equivalent to $\lambda x. \mathbf{height}(x) > \mathbf{height}(\mathbf{m})$

Even though this analysis based on weak-IE can derive the correct truth conditions of plain comparatives, it faces a problem of over-generation. Given the definition of *pos* in (21a), it is in principle possible for *John is taller than Mary* to also have a relative interpretation (23), which says that the height difference between John and Mary “stands out” in context.

- (23) $\llbracket \text{John is } pos \text{ taller than Mary} \rrbracket^C = \mathbf{height}(\mathbf{j}) - \mathbf{height}(\mathbf{m}) \geq \theta_C$

⁸Note that the inequality sign (21c) is a strict one, which is different from the non-strict one in (21a). This is a minor technical problem with Kennedy’s 2007 formulation to unify all three classes of gradable adjectives: for maximum adjectives, the degree is required to be at least the maximum degree (it is impossible to be strictly greater than that), whereas for minimum adjectives, the degree is required to be strictly greater than the minimum degree. Just to be clear, I do not take this to be a serious criticism of Kennedy’s analysis, but I do note that this problem does not arise when we assume that minimum interpretations are in fact derived from a different mechanism than applying *pos*.

In reality, however, this interpretation is never attested. For instance, suppose that Mary is 5'10" and John is a 6'2" basketball player. Given that basketball players are generally quite tall, the 4-inch height difference between John and Mary probably will not stand out among the set of height differences between basketball players and Mary. Even though weak-IE predicts that the relative interpretation is generally dispreferred, without further stipulations, it does allow it to be accessible in certain cases. In particular, the sentences in (24) would be true if the comparative does have a relative interpretation.

- (24) a. # John is not taller than Mary for a basketball player.
 (cf. John is not tall for a basketball player.)
 b. # Compared to other basketball players, John is not taller than Mary.
 (cf. Compared to other basketball players, John is not tall.)
 c. # If we are only talking about basketball players, John is not taller than Mary.
 (cf. If we are only talking about basketball players, John is not tall)

However, these sentences are not true. Moreover, they sound very strange or infelicitous because intuitively it simply does not make sense to try to introduce a comparison class for a plain comparative. We may call such infelicity ungrammaticality, but it is hard to appeal to syntax to explain the infelicity/ungrammaticality. First, even though some analyses treat a *for*-PP as a syntactic argument (e.g., Fults, 2006; Solt, 2011; Bylinina, 2014), they treat it as the argument of *pos* and therefore it is unclear how to block it from appearing in a plain comparative if *pos* is assumed to be there (24a). It seems even more unlikely that the *compared to* phrase in (24b) is an argument rather than an adjunct, so it is not clear why syntactically it cannot appear with plain comparatives. Furthermore, while we can in principle (indirectly) introduce a comparison class by using a conditional, which is clearly an adjunct, it is still infelicitous for plain comparatives (24c). Therefore we probably cannot syntactically rule out the ungrammatical sentences, at least not always.

We probably cannot rule them out pragmatically, either. For instance, we may attempt to appeal to the unnaturalness of the intended meanings. However, the intended meanings are presumably the same as the comparatives modified by *much* (25). Even though the intended meanings are still somewhat unnatural, the sentences (25) sound much better.

- (25) a. John is not much taller than Mary for a basketball player.
 b. Compared to other basketball players, John is not much taller than Mary.
 c. If we are only talking about basketball players, John is not much taller than Mary.

Therefore, I conclude that plain comparatives are semantically incompatible with comparison classes and do not have relative interpretations, whereas the sentences in (25) are fine because the positive form *much* allows for the specification of a comparison class.

The contrast between plain comparatives and positive forms is more striking when we consider *profitable*. Given that Company C made zero profit, the > 0 interpretation of *profitable* and the plain comparative *more profitable than C* are truth-conditionally equivalent, and they are assumed to be derived from the same scale structure. Nevertheless, unlike positive forms, plain comparatives never have relative interpretations (26).

- (26) a. # D was more profitable than C, but not more profitable than C for an oil company.
 (cf. D was profitable, but not profitable for an oil company.)

- b. # Compared to other oil companies, D was not more profitable than C.
(cf. Compared to other oil companies, D was not profitable.)
- c. # If we are only talking about oil companies, D was not more profitable than C.
(cf. Compared to other oil companies, D was not profitable.)

This contrast makes it difficult to use a single mechanism of *pos* to account for *all* the possible interpretations of positive forms and comparatives: either we would under-generate relative interpretations for positive forms with the original IE, or we would over-generate relative interpretations for comparatives with weak-IE. While it is in principle possible for proponents of IE to stipulate additional constraints so that *pos* cannot introduce relative standards for plain comparatives, I suggest that it is worth considering a more straightforward alternative, i.e., *pos* is syntactically incompatible with plain comparatives, which is why they simply lack relative interpretations, and their > 0 interpretations are instead derived by an independent compositional mechanism, which also derives the > 0 interpretations of positive forms. The proposal that *pos* is syntactically incompatible with comparatives can be motivated by the observation that *very* is similarly incompatible with comparatives. Note that *very* can modify a GA and introduces a relative standard (27a).⁹ In contrast, *very* cannot modify a comparative (27b). Presumably, this is not because there is anything wrong with the intended meaning, as it is expressible using *much* (27c). This motivates a syntactic analysis of the ungrammaticality of (27b).

- (27) a. Company E was very profitable _{$\geq \theta$} .
 b. * Company E was very more profitable than C.
 c. Company E was much more profitable than C.

For instance, according to Svenonius and Kennedy (2006), the comparative morpheme adds an additional functional layer to an AP and turns it into a QP, and modifiers such as *very* and *how* can only be combined with APs but not QPs. Given this, if we further assume that *pos* has the same syntactic category as *very*, then it straightforwardly follows that plain comparatives do not have relative interpretations.

4. A compositional analysis of the > 0 interpretation

Now the only remaining issue is how > 0 interpretations of plain comparatives (as well as positive forms) are compositionally derived. In this section, I propose such an analysis.

Recall that Svenonius and Kennedy (2006) and Sawada and Grano (2011) analyze the comparative morpheme as a scale modifier, and in the weak-IE variant discussed in the previous section, *taller than Mary* simply denotes a measure function that returns the height difference between its argument and Mary (22a), which I repeat below in (28).

$$(28) \quad \llbracket \text{taller than Mary} \rrbracket = \lambda x. \text{height}(x) - \text{height}(\mathbf{m})$$

To derive the > 0 interpretation, I propose that instead of applying *pos* to (28), we apply a silent $\emptyset_{\text{SOME}}$ operator which transforms a measure function to an individual predicate saying that the measure of the individual is greater than 0 (29a). We can easily verify that applying $\emptyset_{\text{SOME}}$ to the comparative (28) indeed derives the correct truth conditions as a > 0 interpretation.

⁹This standard is generally intensified, i.e., higher than that of the corresponding positive form, except for maximum-standard adjectives. For instance, *the theater was very full* can be weaker than *the theater was full*.

- (29) a. $\llbracket \emptyset_{\text{SOME}} \rrbracket = \lambda \mu \lambda x. \mu(x) > 0$
 b. $\llbracket \emptyset_{\text{SOME}} \text{ taller than Mary} \rrbracket = \lambda x. \text{height}(x) - \text{height}(\text{m}) > 0$
 c. $\llbracket \text{John is } \emptyset_{\text{SOME}} \text{ taller than Mary} \rrbracket = \text{height}(\text{j}) - \text{height}(\text{m}) > 0$,
 which is equivalent to $\text{height}(\text{j}) > \text{height}(\text{m})$

This proposal is inspired by Schwarzschild and Wilkinson (2002), according to whom the measure phrase in a differential comparative measures the length of the degree interval starting from the degree of the standard to degree of the subject (30b).

- (30) a. $\llbracket \text{two inches} \rrbracket = \lambda J. \mu(J) \geq 2''$
 b. $\llbracket \text{John is } __ \text{ taller than Mary} \rrbracket = [\text{height}(\text{m}), \text{height}(\text{j})]$
 c. $\llbracket \text{John is two inches taller than Mary} \rrbracket = \mu([\text{height}(\text{m}), \text{height}(\text{j})]) \geq 2''$

When there is no overt measure phrase in the comparative, Schwarzschild and Wilkinson (2002) assume that a silent SOME, as defined in (31), is applied to the degree interval.

- (31) $\text{SOME}(J)$ is true iff $\mu(J) \geq \delta$, where δ is determined by context

According to Schwarzschild and Wilkinson (2002), this definition is motivated by the mass quantifier *some*, and they use the contrast in (32) to motivate the introduction of the contextual parameter δ , i.e., whether a very tiny piece of wood is significant or relevant enough to be taken into account depends on the context.

- (32) a. There is some wood in my eye.
 b. There is some wood in my truck.

Essentially, the δ parameter is intended to capture the contextual level of imprecision/granularity. Therefore, if we factor out imprecision/granularity, we can remove the δ parameter from the definition of SOME and use (33) instead.

- (33) $\text{SOME}(J)$ is true iff $\mu(J) > 0$

Also, note that even though Schwarzschild and Wilkinson (2002) use the overt quantifier *some* to motivate SOME, bare mass nouns perhaps provide an even better motivation for this silent operator (34). There is no overt mass quantifier in (34), but the sentence is interpreted in parallel with (32). Therefore, it is reasonable to assume that there is a silent counterpart of *some*, written as $\emptyset_{\text{SOME}}$, that has the same semantics as defined in (33).

- (34) There is water in the glass.

Finally, as Sassoon (2010) points out, according to measurement theory (e.g., Krantz et al., 1971), if degree intervals can be coherently measured, the underlying scale must be an interval scale, in which case the measure of a degree interval simply amounts to the degree difference between its endpoints. Given this, it is easy to see that (29a) and (33) are in fact equivalent for deriving the > 0 interpretation of a plain comparative.

Moreover, this silent operator $\emptyset_{\text{SOME}}$ can straightforwardly derive the > 0 interpretation of *profitable* (35).

- (35) a. $\llbracket \text{profitable} \rrbracket = \text{profit}$

- b. $\llbracket \emptyset_{\text{SOME}} \text{profitable} \rrbracket = \lambda x. \mathbf{profit}(x) > 0$
- c. $\llbracket \text{Company D was } \emptyset_{\text{SOME}} \text{profitable} \rrbracket = \mathbf{profit}(\mathbf{d}) > 0$

The relative interpretation of *profitable* is derived by applying *pos* instead (36). Crucially, since *profitable* has an open scale, existing theories of *pos* all predict that *pos profitable* has a vague relative standard. Therefore, one could retain their favorite analysis of *pos*, be it a probabilistic model or Kennedy's (2007) original formulation, and correctly account for the relative interpretation of *profitable*.

- (36) a. $\llbracket \text{pos profitable} \rrbracket = \lambda x. \mathbf{profit}(x) \geq \theta_C$
 b. $\llbracket \text{Company D was pos profitable} \rrbracket = \mathbf{profit}(\mathbf{d}) \geq \theta_C$

In sum, the proposed analysis provides a new and unified compositional mechanism that derives the > 0 interpretation of a plain comparative, as well as the > 0 interpretation of a positive form such as *profitable*. Together with existing analyses of *pos*, we now have a better understanding of the connections between a zero-standard GA such as *profitable* and a comparative. On the one hand, they both have a scale with a zero degree, which is why they both have a non-vague > 0 interpretation after composing with $\emptyset_{\text{SOME}}$. On the other hand, since only the positive form is syntactically compatible with *pos*, only the positive form can have both a > 0 interpretation and a relative interpretation, whereas the comparative only has a > 0 interpretation.

5. General discussion

5.1. More cases of zero-standard adjectives without minimum degrees

So far, I have been using *profitable* as a case study to show that positive forms in fact can have two types of interpretations and that they are derived from different mechanisms. In particular, since *profitable* patterns with run-of-the-mill minimum adjectives in many respects but its scale structure arguably does not have a minimum degree, I conclude that minimum-standard adjectives are better characterized as zero-standard adjectives. Below, I discuss more examples of zero-standard adjectives whose scale structures do not have a minimum degree. However, I acknowledge that each individual example may not provide evidence as decisive as *profitable*. Depending on one's initial theoretical preference, one might be tempted to explain away with some auxiliary assumptions, and this can be done for most of the examples. However, what I hope to illustrate is that the move from minimum-standards to zero-standards provides a simple explanation that has the best overall coverage of the empirical data, without the need for further stipulations (even if such stipulations can be defended for some individual cases).

Such examples include *early/late*, *fast/slow* (for clocks and watches), and *sharp/flat* (for music instruments), which are discussed by Kennedy (2001). We have already seen in the previous section that there are some reasons to think that *early* and *late* have a full-range scale. Similarly, the examples below provide evidence to think that *fast/slow* and *sharp/flat* have full-range scales (37).

- (37) a. My watch is 2 minutes fast. This clock is 1 minute slow. Therefore, my watch is 3 minutes faster than this clock.
 b. The A string on this guitar is two semitones sharp. The one on that guitar is one semitone flat. Therefore, The A string on this guitar is three semitones sharper than the one on that guitar.

However, these examples are different from *profitable* in that the zero standards in their interpretations are more contextual. For example, the interpretations of *fast* and *slow* make reference to a zero point only when we are talking about instruments that measure time, where there is a standard clock whose speed serves as the zero point. Therefore, one might try to derive the > 0 interpretations of these cases in a special way, while assuming a single scale structure. For example, according to Kennedy (2001), the scale structure of *fast/slow* is always a full range from the absolute zero (i.e., motionless) to infinity, regardless of whether we are talking about clocks or cars. In general, *fast* maps an individual to the interval on the scale from 0 to the individual's speed, and *slow* maps an individual to the interval from the individual's speed to infinity. The interpretations of *fast* and *slow* in (37a) are special in that they map their argument to intervals that extend from the “on time” point to the actual speed, and that they include an additional ZERO function as part of their meanings to shift such intervals to intervals that start from the absolute zero. This is so that they are comparable with measure phrases, which Kennedy (2001) assume denote an interval starting from the absolute zero. Formally, we have semantic representations such as in (38).

- (38) $\llbracket \text{My watch is 2 minutes fast} \rrbracket = \text{ZERO}(\mathbf{fast}_\delta(\mathbf{w})) \succeq \llbracket 2 \text{ minutes} \rrbracket$,
which is equivalent to $\text{ZERO}(\mathbf{fast}_\delta(\mathbf{w})) \supseteq [0, 2\text{min}]$

Even though Kennedy (2001) uses the name ZERO and talks about zero points of the scale in various places in the prose, he actually defines the function ZERO in terms of the minimal element of the scale, presumably under the assumption that the intuitive notion of a zero point can be formally characterized as the minimal element of the scale.¹⁰ As a result, in order for the ZERO function to be defined, the scale structure needs to have a minimal element. While this is the case for *fast/slow* and *sharp/flat*, it is not obviously true for *early/late*. Sure, it is not uncommon for people to hold mythological, religious, or even scientific beliefs that there is a starting point of time, but it seems a little too strong to conclude from this that the scale of *early/late* has a minimal element. Admittedly, people may well be using *early/late* in a way as if the scale had a minimal element, regardless of whether or not they actually believe there is a starting point of time. However, if there is an analysis that has the same or even better empirical coverage and does not have to rely on this specific stipulation about the scale structure of *early/late*, then such an analysis should be preferred.

I suggest that we can have a better analysis by simply treating the zero point as a primitive, instead of characterizing it in terms of the minimal element of a scale. That is, when we specify a scale, we can optionally specify a degree that is considered to be the zero point of that scale. In many cases, such a zero point is lexical and based purely on the property being measured, e.g., the zero point for *tall* is the degree that corresponds to zero amount in height, and the zero point for *profitable* is the degree that corresponds to zero amount of profit. In other cases, the zero point can be more contextual and based on our knowledge about the relevant conventional standard. Here the convention can be widely shared and stable over time, e.g., how fast an accurate clock is supposed to be, but it can also be extremely local and

¹⁰There is a further complication. Note that except when we are talking about instruments that measure time, *fast* is a relative adjective. Given the strict correlation between scale structure and interpretation of the positive form assumed by Kennedy and McNally (2005) and Kennedy (2007), they would probably assume that the scale of *fast* does not have a minimal element, but rather a greatest lower bound. However, whether one tries to characterize the zero point in terms of a minimal element or a greatest lower bound would not affect the discussion below.

highly variable, e.g., the time you and your friend are supposed to meet on various occasions.¹¹ Different adjectives (or different interpretations of an adjective) differ in whether they allow for a scale with a contextual zero point, which can be implemented in terms of semantic selection rules. Under this analysis, the clear-cut interpretations of *profitable*, *early/late*, *fast/slow*, and *sharp/flat* are all instances of the > 0 interpretation, with the only difference being whether the zero point is lexically determined (*profitable*) or contextually determined (the other ones). This analysis does not make more stipulations than Kennedy's. Whenever the interpretation of an adjective is assumed to include ZERO as part of the meaning in Kennedy's analysis, this analysis assumes that it can select for a scale that has a corresponding contextual zero point. Crucially, this analysis does not need to stipulate that the scale of *early/late* has a minimal element, since any degree can in principle be the contextual zero point. In addition, this analysis has better empirical coverage than Kennedy's (2001), which does not and in fact cannot account for *profitable* (since its zero point cannot be characterized as a minimal element). Therefore I conclude that this analysis should be preferred to Kennedy's (2001).

5.2. Silent morphemes and a potential worry of over-generation

According to my analysis, there are two silent morphemes responsible for generating the interpretation of a positive form. One is *pos*, which is assumed to be a silent counterpart of *very* and is responsible for the relative interpretation. The other is $\emptyset_{\text{SOME}}$, which generates the > 0 interpretation.

So far I have not imposed any restrictions on the distribution of $\emptyset_{\text{SOME}}$. This means that a positive form whose scale structure has a zero degree is predicted to have a > 0 interpretation.

In some cases, this does not cause a major problem. For instance, given that **height** has a zero degree (even though it is not realized by any individual), the current analysis predicts that *John is tall* can have a > 0 interpretation, i.e., **height**(j) > 0 . While this interpretation is not attested, we can presumably rule it out by noting that it is a tautology, which means that it is totally uninformative and therefore is almost never picked up by pragmatic language users.

However, this explanation is not always applicable. For instance, given that **fullness** also has a zero degree, the current analysis predicts that *full* has a > 0 interpretation, which would mean *non-empty*. Since *this glass is non-empty* is not a tautology, we cannot say that the predicted > 0 interpretation of *full* is not attested because it is totally uninformative. Similar cases exist for relative adjectives as well. For instance, while **likelihood** has a zero degree, *likely* does not have a > 0 interpretation, which would just mean *possible*.

Therefore, it seems that the current analysis faces a problem of over-generation, which cannot always be explained by pragmatics. Below I discuss another solution, but I acknowledge that it is not completely satisfying, either.

The idea is to assume that $\emptyset_{\text{SOME}}$ has further selectional requirements. Recall our earlier example *there is water in the glass*, which motivates our definition of $\emptyset_{\text{SOME}}$. Now consider the

¹¹Note that a zero point is contextual iff it cannot be determined based on the property being measured (i.e., it is not lexical) and therefore needs to be supplied by context. This definition does not entail that the zero point always has to vary in different contexts. In fact, when the relevant convention is widely shared and stable over time (e.g., the standard speed of a clock/watch), the zero point can be highly stable across contexts.

ungrammatical sentence **there is student in the classroom*. It seems reasonable to assume that the ungrammaticality is due to a syntactic restriction according to which a bare singular count noun cannot have an existential interpretation. If we adopt a similar strategy, we would say that the syntactic distribution of $\emptyset_{\text{SOME}}$ is further restricted so that it does not combine with *full* or *likely* for pure syntactic reasons.

Of course, in order for this solution to be more than a pure stipulation, we should look for morpho-syntactic evidence that helps predict whether the > 0 interpretation is available. To some extent, we can indeed find such evidence. There are morphemes that are good indicators of the > 0 interpretation (39).

- (39) a. *-y*: *salty, spicy, dirty, windy, ...*
 b. *-ive*: *effective, supportive, active...*
 c. *-ful*: *helpful, harmful, colorful ...*
 d. *-able*: *profitable, curable, noticeable, accessible, ...*
 e. *-ed*: *spotted, striped, bent, curved...*

The list is certainly not exhaustive, but it does seem to provide some hope that maybe we can restrict the syntactic distribution of $\emptyset_{\text{SOME}}$ to rule out the over-generation cases. However, this solution is not entirely satisfying, either, because these morphemes are not completely reliable indicators of the > 0 interpretation, and there can still be an over-generation problem. For instance, *expensive/pricy/closed*, despite having the relevant morphemes, do not seem to have a > 0 interpretation, which would mean *non-free/not fully open*.

That said, I should note that previous accounts of gradable adjectives also face this problem of over-generation. An IE-based analysis says nothing about whether a gradable adjective with a fully closed scale would use the maximum or minimum degree as the standard and therefore would have to stipulate what counts as a natural transition. Probabilistic models can generate the > 0 interpretation for some probability distributions. Setting aside the problem that they wrongly predict the > 0 interpretation to be vague for *profitable* (which does not arise for closed-scale adjectives), since they do not have a full theory of how the probability distribution is determined, they do not have a principled way to avoid over-generating the > 0 interpretation, either. Therefore, determining whether an adjective has a > 0 interpretation is still an open problem, which I will leave for future research.

6. Conclusion

In this paper, I argued that the so-called minimum-standard gradable adjectives should in fact be characterized as having zero standards, which may not be minimum degrees. Moreover, I argued that such > 0 interpretations should be derived by an independent compositional mechanism, i.e., by applying a silent $\emptyset_{\text{SOME}}$ to the scale. Together with existing treatments of *pos*, the proposed ambiguity analysis provides a better understanding of the similarities and differences between positive forms and comparatives. However, it remains an open issue how to restrict the distribution of $\emptyset_{\text{SOME}}$ to avoid the problem of over-generating > 0 interpretations.

References

- Bierwisch, M. (1989). The semantics of gradation. In M. Bierwisch and E. Lang (Eds.), *Dimensional Adjectives: Grammatical Structure and Conceptual Interpretation*, pp. 163–192. Berlin Heidelberg: Springer-Verlag.
- Bylinina, L. (2014). *The Grammar of Standards: Judge-dependence, Purpose-relativity, and Comparison Classes in Degree Constructions*. Phd thesis, Utrecht University.
- Fulst, S. (2006). *The structure of comparison: an investigation of gradable adjectives*. Phd thesis, University of Maryland.
- Kennedy, C. (2001). Polar opposition and the ontology of ‘degrees’. *Linguistics and Philosophy* 24, 33–70.
- Kennedy, C. (2007). Vagueness and grammar: The semantics of relative and absolute gradable adjectives. *Linguistics and Philosophy* 30, 1–45.
- Kennedy, C. and L. McNally (2005). Scale structure, degree modification, and the semantics of gradable predicates. *Language* 81(2), 345–381.
- Krantz, D. H., R. D. Luce, P. Suppes, and A. Tversky (1971). *Foundations of measurement: Additive and Polynomial Representations*, Volume 1. New York: Academic Press.
- Lassiter, D. (2017). *Graded Modality: Qualitative and Quantitative Perspectives*. Oxford University Press.
- Lassiter, D. and N. D. Goodman (2013). Context, scale structure, and statistics in the interpretation of positive-form adjectives. In T. Snider (Ed.), *Semantics and Linguistic Theory (SALT)* 23, pp. 587–610. LSA and CLC Publications.
- Lassiter, D. and N. D. Goodman (2015). Adjectival vagueness in a Bayesian model of interpretation. *Synthese*, 1–36.
- Qing, C. and M. Franke (2014). Gradable adjectives, vagueness, and optimal language use: A speaker-oriented model. In T. Snider, S. D’Antonio, and M. Wiegand (Eds.), *Semantics and Linguistic Theory (SALT)* 24, Ithaca, NY, pp. 23–41. CLC.
- Sassoon, G. W. (2010). Measurement theory in linguistics. *Synthese* 174(1), 151–180.
- Sawada, O. and T. Grano (2011). Scale structure, coercion, and the interpretation of measure phrases in Japanese. *Natural Language Semantics* 19(2), 191–226.
- Schwarzschild, R. (2005). Measure phrases as modifiers of adjectives. *Recherches linguistiques de Vincennes* 30(2), 207–228.
- Schwarzschild, R. and K. Wilkinson (2002). Quantifiers in comparatives: A semantics of degree based on intervals. *Natural Language Semantics* 10(1), 1–41.
- Solt, S. (2011). Notes on the comparison class. In R. Nouwen, R. van Rooij, U. Sauerland, and H.-C. Schmitz (Eds.), *Vagueness in Communication*, pp. 127–150. Springer.
- Solt, S. (2012). Comparison to arbitrary standards. In A. Aguilar-Guevara, A. Chernilovskaya, and R. Nouwen (Eds.), *Sinn und Bedeutung (SuB)*, Volume 16, pp. 557–570.
- Svenonius, P. and C. Kennedy (2006). Northern Norwegian degree questions and the syntax of measurement. In M. Frascarelli (Ed.), *Phases of interpretation*, pp. 129–157. The Hague: Mouton de Gruyter.