

Evaluativity: a proposal for empirical investigation

Floris T. van Vugt
florisvanvugt@ucla.edu

June 7, 2010

1 Markedness

Adjectives that denote into scales, such as *long/short*, often come in pairs that are asymmetric. For example, Clark (1969) observes that only one can denote the scale itself (e.g. *length* but **shortness*), and only one can combine with a measure phrase, as in (1). Traditionally, the adjective that has the more restricted distribution (e.g. *short*) has been called *marked*, and its contrary *unmarked*. Similarly, it is felt that when the marked term is used in questions this is understood as presupposing that it applies to the argument. For example, asking (2-b) seems to presuppose that John is short, which contrasts with (2-a) which seems to presuppose nothing about John, not even that he would be tall. Finally, it is reported that marked terms are learned later during child language acquisition (Clark, 1972).

(1) a. John is 5ft tall.

14 b. *John is 5ft short.

15 (2) a. How tall is John?

16 b. How short is John?

17 In the seventies a number of experimental studies was performed that consti-
18 tute evidence for the psychological reality of the marked–unmarked distinc-
19 tion (see for example Seymour (1974)). Chase & Clark (1971) investigated
20 the marked/unmarked pair *below/above* and report that subjects had more
21 difficulty affirming that a star was below the circle than that the circle was
22 above the star. The salient explanation is that *below* denotes a concept that
23 is encoded in a somehow more complex way than *above* (for a more detailed
24 discussion of these sentence–picture verification tasks, see Chase & Clark
25 (1972)).

26 In a recent criticism, Proctor & Cho (2006) suggests that these experi-
27 mental results can be explained in a more general framework, in which the
28 advantage to affirming the *above*–sentences relative to affirming the *below*–
29 sentences stems from the fact that the affirmative response itself has some
30 abstract positive polarity. This positive polarity aligns with the stipulated
31 polarity *above* but will produce a mismatch with that of *below*, causing the
32 reaction time differences (a version of this idea has also been presented in
33 Carpenter & Just (1975)). Though a detailed survey of this discussion is
34 beyond the scope of this paper, it is important to note that the reaction time
35 difference that was reported in the seventies may not be due to processing

36 difficulty of marked terms in themselves, but rather their interplay with the
37 required response.

38 1.1 Markedness in comparatives

39 Higgins (1977) reported that comparatives with marked terms are more pre-
40 suppositional than those with unmarked terms. Presuppositionality is un-
41 derstood as follows. When someone utters (3-a), in some accounts this pre-
42 supposes that both Bob and Fred are bad. However, (3-b) does not seem to
43 presuppose either Bob and Fred being bad or good.

- 44 (3) a. Bob is worse than Fred
45 b. Bob is better than Fred
46 c. Fred is better than Bob
47 d. Fred is worse than Bob

48 Higgins (1977) measured the presuppositionality in an *acceptability* task,
49 where subjects were asked to rate the acceptability of a sentence that com-
50 pared two items that clearly had a quality opposite to the one implied by the
51 adjective. Such sentences with a marked adjective, e.g. (4-b), were judged
52 less acceptable than those with an unmarked one, e.g. (4-a).

- 53 (4) a. A feather is heavier than a snowflake
54 b. A mountain is lighter than a ship

55 The author argues that both results can be explained by the marked adjectives carrying the presupposition that the entities that are compared possess
56 the marked quality. For example, (3-a) implies that both Bob and Fred are
57 bad, but (3-c) does not imply that they are good, hence they are perceived
58 as less synonymous. Similarly, the use of the marked adjective in (4-b) implies that the arguments are light, which is not the case, causing subjects to
59 perceive the sentence as less acceptable.

62 I would argue, however, that the fact that marked adjectives are much less
63 frequent than unmarked adjectives caused participants to relatively disprefer
64 a sentence with a marked adjective. In order to control for this, it would have
65 been desirable to compare ratings of sentences in (5). If the acceptability
66 difference is absent here, this would constitute evidence that it is due to the
67 adjective markedness and not some other factor.¹

- 68 (5) a. A guilder is heavier than a dollar.
69 b. A guilder is lighter than a dollar.

70 2 Evaluativity

71 In more modern semantic terminology the effect reported in section 1.1 is
72 referred to as *evaluativity*. A phrase is evaluative if “it makes reference to
73 a degree that exceeds a contextually specified standard” (Rett, 2008a). For

¹When I propose use of Higgins (1977)’s experimental paradigm I will assume that these appropriate controls are performed as well.

example, uttering (6-a) establishes that Boris exceeds a contextually specified standard of tallness. However, (6-b) implies no such thing, and similarly (6-c). (6-d) is again commonly perceived as implying that the individuals that are mentioned are short.

- (6)
- a. Boris is tall.
 - b. Boris is taller than Doris.
 - c. Boris is as tall as Doris.
 - d. Boris is as short as Doris.

Rett (2008b) suggests that markedness plays a role in the evaluativity of comparatives and equatives. In particular, she argues that comparatives are not generally evaluative, regardless of whether marked or unmarked terms are used. This property is referred to as *polarity-invariance*. The equative construction with a marked adjective, however, is usually perceived as evaluative (e.g. (6-d)), and hence the equative is *polarity-variant*.

2.1 Evaluativity in the equative

How can this be explained? Let us first consider the equative. (6-c) could be construed to be ambiguous between (7-a) and (7-b). Now (6-d) can be interpreted analogously by (7-c) or (7-d).

- (7)
- a. $\exists d \text{ MAX}\{d | \text{TALL}(\text{Boris}, d)\} = \text{MAX}\{d | \text{TALL}(\text{Doris}, d)\}.$
 - b. $\exists d \text{ MAX}\{d | \text{TALL}(\text{Boris}, d)\} = \text{MAX}\{d | \text{TALL}(\text{Doris}, d)\} > d_{\text{tall}}$ for

- 94 some contextually specified standard d_{tall} .
- 95 c. $\exists d \text{ MAX}\{d|\text{SHORT}(\text{Boris}, d)\} = \text{MAX}\{d|\text{SHORT}(\text{Doris}, d)\}$.
- 96 d. $\exists d \text{ MAX}\{d|\text{SHORT}(\text{Boris}, d)\} = \text{MAX}\{d|\text{SHORT}(\text{Doris}, d)\} > d_{\text{short}}$
- 97 for some contextually specified standard d_{short} .

98 Now the crucial observation is that TALL and SHORT denote onto the
99 same scale, but in opposite directions. The result is that the maximal degree
100 to which a person is tall is automatically the maximal degree to which the
101 person is short². As a consequence, (7-a) and (7-c) are equivalent. Notice
102 that (7-b) and (7-d) are not equivalent since the contextual standards for the
103 long and short scales may well differ.

104 The next step in the reasoning is that since (7-a) and (7-c) are equiva-
105 lent, they enter into semantic competition. This means that in some way
106 they compete for which is the most efficient way of expressing their message.
107 Now (7-c) uses a marked term, contrary to (7-a), and since there is no other
108 difference between them, one can say (7-c) is more marked overall and there-
109 fore dispreferred³. As a result, (7-c) is blocked as a reading of (6-d) since the
110 same message could have been conveyed more efficiently.

111 As a consequence, (7-d) is the only remaining reading, which means that
112 (6-d) is disambiguated and, in the absence of other factors, will always be
113 interpreted evaluatively. Compare, however, with (6-c) which can be evalua-

²Here in the former case “maximal” is understood relative to the canonical ordering on TALL scale, and in the latter relative to the inverse ordering, since SHORT is the antonym of TALL.

³The reason for this is not made explicit in Rett (2008b) but is plausible given earlier accounts of how marked terms are more rare and might take more time to process.

114 tive or not evaluative. Consequently, we cannot deduce from (6-c) that Boris
 115 and Doris are tall, which suffices to classify it as not evaluative.

116 2.2 Evaluativity in the comparative

117 Comparatives with unmarked adjectives such as in (8-a) are generally agreed
 118 upon not to be evaluative. On the other hand, there is disagreement in the
 119 literature as to whether comparatives with marked adjectives, e.g. (8-b), are
 120 evaluative.

- 121 (8) a. Boris is taller than Doris.
 122 b. Boris is shorter than Doris.

123 Clark (1969) writes that “‘Pete is worse than John’ unambiguously impl[ies]
 124 negative evaluations of Pete and John” (p.391). That is, marked compara-
 125 tives are seen as evaluative. However, Rett (2008b) argues that upon closer
 126 scrutiny, comparatives are not evaluative.⁴

127 Indeed, that comparatives are not evaluative follows fairly seamlessly from
 128 the analysis presented before for equatives. Let us assume that (8-b) is
 129 ambiguous between the evaluative and non-evaluative reading in (9-a) and
 130 (9-b).

- 131 (9) a. $\text{MAX}\{d | \text{SHORT}(\text{Boris}, d)\} > \text{MAX}\{d | \text{SHORT}(\text{Doris}, d)\}$

⁴Except, of course, comparatives with extreme adjectives, which are always perceived as evaluative. For example, *Tim is more moronic than Pete* clearly implies a judgement about the intelligence or absence thereof of the individuals in question. For the sake of simplicity, I will exclude these extreme adjectives from our discussion.

- 132 b. $\text{MAX}\{d|\text{SHORT}(\text{Boris}, d)\} > \text{MAX}\{d|\text{SHORT}(\text{Doris}, d)\} > d_{\text{short}}$
- 133 c. $\text{MAX}\{d|\text{TALL}(\text{Boris}, d)\} > \text{MAX}\{d|\text{TALL}(\text{Doris}, d)\}$
- 134 d. $\text{MAX}\{d|\text{TALL}(\text{Boris}, d)\} > \text{MAX}\{d|\text{TALL}(\text{Doris}, d)\} > d_{\text{tall}}$

135 Now the non-evaluative reading (9-a) cannot enter into competition with
 136 the reading in (9-c), where the marked adjective is replaced by its unmarked
 137 counterpart. The problem is that they do not mean the same thing, and
 138 therefore they do not enter into semantic competition. Thus, none of the
 139 readings is blocked and as a result, the marked comparative is not evaluative.

140 **2.3 Critique of non-blocking analysis**

141 The analysis presented in section 2.2 is appealing since the ambiguity that is
 142 ascribed to comparatives and evaluatives can explain why there are contexts
 143 in which they are evaluative and others in which they are not. Further-
 144 more, this account is supported by the variability in the presuppositional-
 145 ity observed by Higgins (1977), who remarks that “comparatives containing
 146 marked adjectives from a ratio scale *can* be interpreted neutrally”⁵.

147 However, the same studies’ finding that marked comparatives are in gen-
 148 eral more presuppositional is not in line with the analysis. If we are to
 149 interpret this lack of experimental confirmation to problems in its design,
 150 then we will arguably also lose its support for Rett (2008b)’s analysis of
 151 comparatives.

⁵Emphasis added. The author defines ratio adjectives as those that can combine with a measure phrase and that have a clear zero point.

152 Also, the argument for the non-evaluativity of marked comparatives feels
 153 somewhat unsatisfying. The crucial step was to compare the reading (9-a)
 154 with (9-c). But the latter seems a rather surprising choice as competitor for
 155 (9-a). What we essentially have done is taken (10-a) and compared it with
 156 (10-b), concluding that they are not synonymous. On what grounds was
 157 *taller* even considered as a candidate? Notice that in general a sentence with
 158 *smaller* implies the *negation* of the same sentence with *larger*, so it seemed
 159 we could not have chosen a worse candidate for equivalence. And what is
 160 more, why is the synonymous (10-c) excluded as a candidate?

- 161 (10) a. Boris is shorter than Doris (non-evaluative)
 162 b. Boris is taller than Doris (non-evaluative)
 163 c. Doris is shorter than Doris (non-evaluative)
 164 d. Boris is not taller than Doris (non-evaluative)

165 Rett (2008a) observes that apparently the switching of the arguments has
 166 blocked the semantic competition. Interestingly, a similar result might be
 167 derived from the *principle of the primacy of functional relations* (Clark, 1969).
 168 Or, perhaps a less strong restriction could be that pairs can enter in semantic
 169 competition only if they differ minimally, where *minimal difference* could be
 170 defined as a relation between sentences α and β that hold if (i) $\alpha \neq \beta$, and
 171 (ii) there is no sentence γ that is less different from α than β is⁶ and that
 172 occurs at some point in a stepwise transformation from α to β .

⁶Of course some distance metric is implicit here. It could be a sort of Levenshtein distance on strings of words.

173 3 Experimental investigation

174 I will argue here that the proposed analysis of evaluativity needs to be
175 founded on a more firm experimental investigation, so that our theories are
176 informed not only by the intuition of those who design them, but also by
177 more objective data revealing how people use the sentences in question.

178 3.1 Comparing comparatives and equatives: a first ex- 179 perimental proposal

180 For example, to the best of my knowledge, a presuppositional analysis such
181 that of Higgins (1977) has not been performed for equatives. Higgins inves-
182 tigated various types of comparatives to see how much presupposition they
183 carried relative to each other. In order to test the theory that has been
184 presented here it will be crucial to gain insight into how presuppositional
185 equatives are relative to comparatives. Rett (2008b) predicts that they are
186 much stronger in what they presuppose. This can be tested by a paradigm
187 adapted from Higgins (1977).

188 We present subjects an *acceptability* task. We make a list of pairs of
189 non-extreme adjectives, one of which is marked and the other one not. For
190 both adjectives in the pair we find two objects who clearly do not possess the
191 denoted property⁷. For example, for the *tall-short* pair, we could take *dwarf*,
192 *miniature* as candidates for (not) *tall* and *skyscraper*, *poplar* for (not) *short*.

⁷To ensure comparability with the Higgins (1977) study, one can copy the examples used.

193 We present subjects with sentences of the form “X is as A as Y,” where A is
194 an adjective and X and Y the candidates that clearly do not have property
195 A. Subjects are then asked to rate the acceptability by clicking with a mouse
196 somewhere on a bar ranging from 0 for totally unacceptable to 1 for totally
197 acceptable.

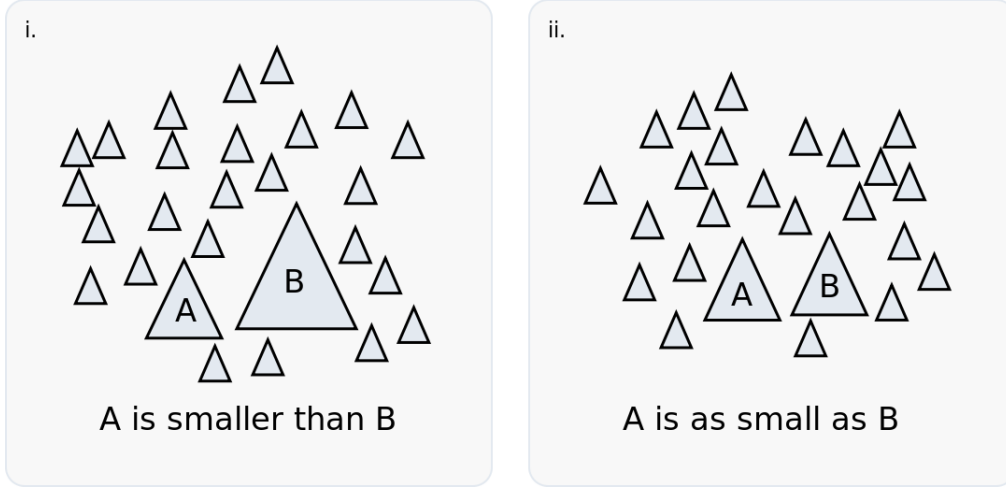
198 In addition to these we test the subjects on the marked–unmarked com-
199 parative from Higgins (1977)’s original study in order to ensure we replicate
200 the effect and in order to provide a benchmark for the effect size of the
201 equative.

202 Our theory predicts that the difference in acceptability between this equa-
203 tive marked–unmarked pair will be greater than that between the compara-
204 tive marked–unmarked.

205 **3.2 Context–sensitivity of comparatives and equatives**

206 The problem in a Higgins (1977)–like approach to presuppositionality in com-
207 paratives and equatives is that we rely on subject’s judgements independent
208 of any context. This means that it is possible that the task becomes met-
209 alinguistic and therefore sensitive to many factors that come into play when
210 people are asked to freely reflect on their opinion. For example, people might
211 try to come up with a context or natural communication setting in which
212 certain readings are appropriate, and thus their response would be a measure
213 of their creativity much more than anything else. It would be preferable to
214 address the issue of presuppositionality in a more direct way by making up

Figure 1: Equative and comparative embedded in context

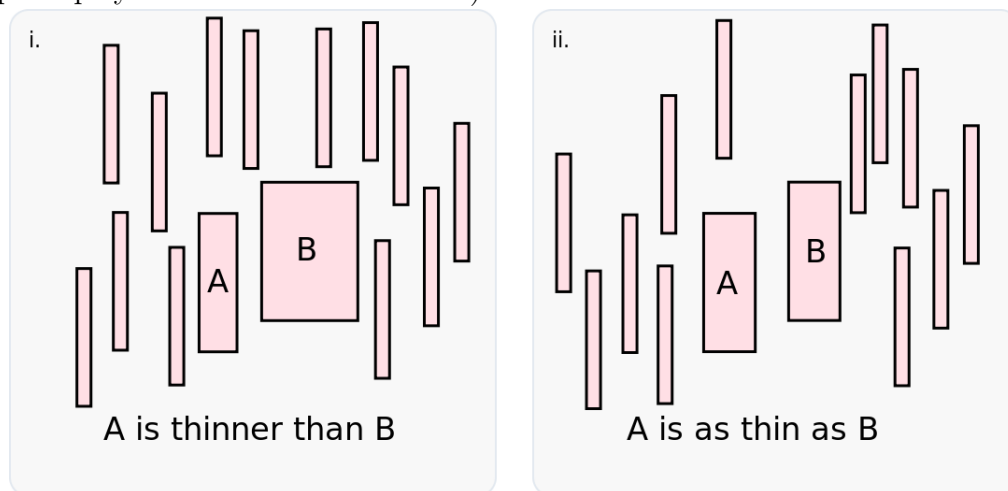


215 a concrete situation in which the judgements of people can be compared.

216 I propose an experiment in which a context is provided for two objects
 217 A and B that are compared for size by placing them in a field of smaller
 218 items. This means they are both relatively large. If our theory is correct,
 219 then that means that the equative *A is as small as B* will be dispreferred as
 220 a description when they are equal in size, since both are not small. However,
 221 when they differ in size, then *A is smaller than B* should be fine, since we
 222 can interpret it non-evaluatively and in that case it will be true. This is
 223 illustrated in figure 1 where the reader is invited to introspectively verify his
 224 own acceptability judgements.

225 A first part of this experimental program would be a pilot study where
 226 these pictures are given to subjects who are asked to rate them on a continu-
 227 ous scale. We predict that this will yield the same result as the acceptability

Figure 2: Using different adjective pairs to test the same predictions (or perhaps yield a different intuition?)



judgement task from the previous section, there the equative is significantly
less acceptable than the comparative.⁸ In order to make the purpose of the
task less obvious to the participant, it will be sensible to include also the
same cases but with a context of large objects. This will furthermore pro-
vide a baseline response against which the acceptability judgements of the
two crucial cases can be compared. Also, the experiment can be peppered
with other adjectives for which similar comparative and equative pictures
can be drawn, for instance as shown in figure 2.

⁸I verified this informally with a naive subject who told me he hesitated tremendously to call the equative correct in the case of equating large objects in a small context by using *as small as*.

236 3.3 Picture–production paradigm

237 Once the results from this pilot study are established, we can move on to a
238 more complex task in which we will simulate production by allowing the par-
239 ticipant to choose from different utterance options which one best describes
240 the picture in question.

241 In figure 3 the stimuli for the experiment are shown. Let us first consider
242 the case of the equatives. We expect that in a LARGE context, both *A is as*
243 *small as B* and *A is as large as B* are possible descriptions, since the latter
244 can be interpreted non-evaluatively. In a SMALL context, *A is as small as B*
245 is predicted to be not possible as a description since it can only be interpreted
246 evaluatively, and *A* and *B* are not small, but large. This should be reflected
247 in the overall participant’s choice pattern.

248 Now in the case of the comparatives there are two possible answer schemas.
249 Take the example of *A* being smaller than *B*. One schema (the *smaller*
250 schema, cf. figure 3) proposes a choice between *A is smaller than B* and
251 *A is larger than B*. These are the two sentences that are candidates for
252 semantic competition in Rett (2008b). Notice that the latter is false; there-
253 fore all participants should choose the former if they are performing the task
254 correctly.

255 In a second answer schema, referred to as *invert*, however, the participant
256 can choose between *A is smaller than B* and *B is larger than A*. In this
257 case, both answers are true in their logical sense. Rett (2008b) suggests that
258 neither is presuppositional, and therefore neither is excluded for that reason.

259 This means that we expect to see no difference in choice pattern between
 260 these phrases in the LARGE context, nor in the SMALL context. If, however,
 261 the switching of the arguments is not as fundamentally disruptive as has been
 262 assumed, then we expect a preference for the use of the unmarked term in
 263 both contexts since apart from markedness of the term and the order of the
 264 arguments the utterances are identical.⁹ Furthermore, reaction times might
 265 provide a clue as to the perceived difficulty or hesitation of the participants.

266 3.4 Time–course analysis of semantic competition

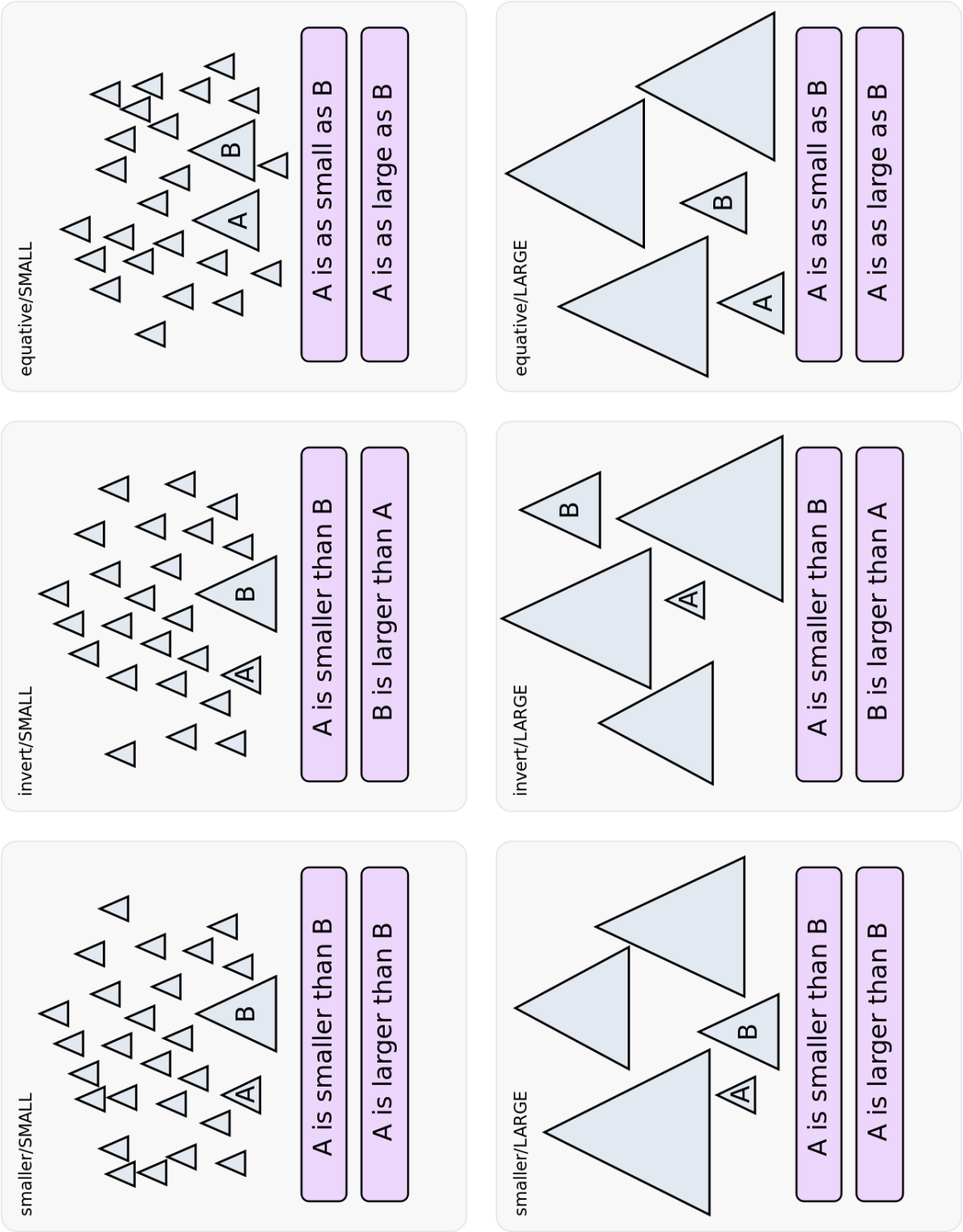
267 The semantic competition account provides a further possibility for exper-
 268 imental verification. The competition is in an abstract way comparable to
 269 the way Gricean implicatures are computed by a listener. Such implicatures
 270 are calculated as follows. If a listener hears a sentence ϕ and then consid-
 271 ers a logically stronger sentence ψ that would have taken the same effort to
 272 produce, then he or she will conclude that the speaker thinks ψ is false. For
 273 otherwise, the speaker would have uttered ψ to be maximally informative.

274 If we assume for a moment that the speaker is intending to say that two
 275 objects A and B are equal in vertical size. Then he or she considers uttering
 276 one of (11). That is, the two are in competition. Now suppose that there is
 277 a Gricean–like maxim that dictates: *say what you have to say as efficiently*
 278 *as possible*, briefly: *be efficient*¹⁰. Now since (11) mean the same thing and

⁹The appeal of this experiment lies precisely in the comparison between the contexts in this case to be highly informative with respect to our theories.

¹⁰Perhaps this can be seen as a special case of the *maxim of manner* that requires us

Figure 3: Experimental design for the picture-production task



279 therefore convey exactly the same information, the usage of *short* is less
280 efficient than *tall* since it is more marked. This means that the speaker will
281 utter (11).

- 282 (11) a. *A* is as tall as *B*.
283 b. *A* is as short as *B*.

284 At this point, one should remark that nothing in the theory of semantic
285 competition has committed us to this view that the competition unfolds in
286 real time while the subject is preparing the utterance. This is analogous to
287 how the theory of Gricean pragmatics does not imply that this implicature
288 is calculated every time by the subject. For all we know it could also be
289 hard-wired into the meaning of the word.

290 However, in the case of pragmatic implicatures Bott & Noveck (2004)
291 showed that subjects who were told that *some* means “some or possibly all”,
292 i.e. the logical meaning of *some*, responded faster to verification studies than
293 a different group of subjects who were instructed that it meant “some but
294 not all”, i.e. the pragmatic meaning. Also, subjects who were not instructed
295 any particular meaning for *some*, responded according to the logical meaning
296 more often when they were put under time pressure to respond. The authors
297 conclude that calculating the pragmatic implicature takes time and that it
298 is derived “on-line” every time the word *some* is used.

to be as clear as possible.

299 3.5 An experimental proposal for competition annihi- 300 lation

301 This means that it is possible, though by no means necessary, that the se-
302 mantic competition happens in real time. In this case we would be able to
303 make people use the marked equative non-evaluatively.

304 The data from the experiment described in section 3.3 is needed for our
305 first step. We investigate at what latencies subjects respond. Now a STRICT
306 time limit is decided so that exactly 50% of the responses of the pilot subjects
307 fall before and the rest after this time limit. Further, a LONG time limit is
308 decided so that 90% of the responses is included.¹¹ Now the actual test
309 subjects are divided into two groups. One group is given the STRICT time
310 limit, the other the LONG time limit.

311 Our hypothesis that the semantic competition happens in a separate
312 stage, after other picture-encoding decisions are taken, and therefore takes
313 time makes the following prediction. Under the STRICT time limit, the equa-
314 tive in the SMALL context will be equally often described with *smaller*
315 or *larger*, even though the pilot test presumably shows that it is dispreferred
316 to use *smaller* in that context. However, in the LONG time limit, there should
317 be a significant preference for *larger*, i.e. a replication of the results in the
318 previous study without time-limit.

¹¹We on purpose do not include all responses since (a) obviously there will be outliers, but also (b) it is important that subjects have at least some sense of time pressure in both cases, though in one case it is much more severe.

319 The same comparison can be made for the comparative in the *invert* con-
320 dition (cf. figure 3). Depending on what effect we found in the earlier study
321 without time pressure, seeing whether this *invert* condition is affected in the
322 same way as the equative by increased time pressure will allow us to gain
323 insight into the extent to which their evaluativity or not is comparable. Fi-
324 nally, the *smaller* condition (cf. figure 3) serves as a crucial control condition,
325 since one of the examples is strictly wrong. This is vital if we would find that
326 subjects choose equally often either response in the *invert* condition, which
327 could be interpreted as a result of too high time pressure. Only when they
328 do *not* respond at chance level in the *smaller* condition can we rule out this
329 interpretation.

330 4 Conclusion

331 Certain degree scales are denoted into by pairs of opposite adjectives that
332 are asymmetric in that one is the default, unmarked case and the other is
333 its marked alternative. Phrases that relate two objects along a particular
334 domain using such adjectives are often felt to be evaluative in the marked
335 equative construction but not in the marked comparative, nor in any of the
336 constructions using unmarked adjectives. In this paper, several experiments
337 are proposed in full detail that can further clarify how these evaluativity
338 patterns are used by human subjects, so that our finest semantic theories
339 can be informed by rigorous empirical results.

References

- Bott, Lewis, & Noveck, Ira. 2004. Some utterances are underinformative: The onset and time course of scalar inferences. *journal of memory and language*, **51**, 437–457.
- Carpenter, Patricia A., & Just, Marcel Adam. 1975. Sentence comprehension: A psycholinguistic processing model of verification. *Psychological review*, **82**(1), 45–73.
- Chase, W.G., & Clark, H.H. 1971. Semantics in the perception of verticality. *British journal of psychology*, **62**, 211–216.
- Chase, W.G., & Clark, H.H. 1972. On the process of comparing sentences against pictures. *Cognitive psychology*, **3**, 472–517.
- Clark, Eve V. 1972. On the child’s acquisition of antonyms in two semantic fields. *Journal of verbal learning and verbal behavior*, **11**(6), 750 – 758.
- Clark, Herbert H. 1969. Linguistic processes in deductive reasoning. *Psychological review*, **76**(4), 387–404.
- Higgins, E. Tory. 1977. The varying presuppositional nature of comparatives. *Journal of psycholinguistic research*, **6**(3), 203–222.
- Proctor, Robert W, & Cho, Yang Seok. 2006. Polarity correspondence: A general principle for performance of speeded binary classification tasks. *Psychological bulletin*, **132**(3), 416–442.
- Rett, Jessica. 2008a. Antonymy and evaluativity. In: Gibson, M., & Friedman, T. (eds), *Proceedings of salt xvii*. CLC Publications.
- Rett, Jessica. 2008b. *Degree modification in natural language*. Ph.D. thesis, Rutgers University.
- Seymour, Philip H. K. 1974. Stroop interference with response, comparison, and encoding stages in a sentence-picture comparison task. *Memory & cognition*, **2**(1A), 19–26.