

## **UC Merced**

### **Proceedings of the Annual Meeting of the Cognitive Science Society**

#### **Title**

Discriminative learning predicts human recognition of English blend sources

#### **Permalink**

<https://escholarship.org/uc/item/3p38k0gw>

#### **Journal**

Proceedings of the Annual Meeting of the Cognitive Science Society, 36(36)

#### **ISSN**

1069-7977

#### **Authors**

Seyfarth, Scott  
Myslin, Mark

#### **Publication Date**

2014

Peer reviewed

# Discriminative learning predicts human recognition of English blend sources

Scott Seyfarth (sseyfarth@ucsd.edu)

Mark Myslin (mmyslin@ucsd.edu)

Department of Linguistics

University of California, San Diego

9500 Gilman Drive, La Jolla, CA 92093-0108

## Abstract

Strict compositionality in morphological theory is problematic for explaining how language-users comprehend phenomena like the partial yet non-decomposable forms in phonaesthemes and in blends like *edutainment*. An alternative account, based on discriminative learning, proposes that language-users associate linguistic cues (e.g., short segment or letter strings) with multiple simultaneous possible lexical and grammatical meanings. We evaluate this account on off-line human identifications of partial word-forms, using English blend words as our test case. We hypothesize that readers' ability to parse out source meanings from written blend forms should be correlated with how strongly a naïve discriminative reading model associates the cues in each form with the correct source meanings. We provide evidence for this claim in two experiments, in which the discriminative learning model reliably predicted participants' success rate in guessing the sources of both attested and novel blends. This finding supports discriminative learning as a realistic model of how readers parse wordforms and map them to meanings. Further, the result points towards a novel, precise account of blend processing.

**Keywords:** blends; discriminative learning; parsing; morphological processing; reading

## Introduction

Language-users are able to recognize the constituents of morphologically-complex words like *rewrite*. Under the dominant view of word processing, they accomplish this by decomposing words into stems and affixes, and during this process they activate the appropriate lexical and semantic representations for the constituents (Taft & Forster, 1975; Taft, 1981; Stockall & Marantz, 2006). Among other things, this process requires representations for combinative forms, such as affixes, which are most likely acquired by distributional learning over sets of words that share semantic and phonological properties (e.g., Finley & Newport, 2011; Finley & Wiemers, 2013).

However, not all partial forms can be separately mapped to distinct meanings. For example, phonaesthemes and blends like *edutainment* include partial yet non-decomposable forms. While a blend is composed of multiple constituents (*education* + *entertainment*), it is not generated by any regular morphological process. There is no rule for concatenating the first two syllables of one word and the last two syllables of another word to create a wordform with a related meaning. Furthermore, blends are not decomposable in the same way that compounds like *blackboard* are decomposable: neither *edu* nor *tainment* exists as an independent word. Blends are also not decomposable in the way that inflected forms like *walked* or derived forms like *naturalness* are decomposable. For any particular blend, neither of its partial forms likely

exists as a constituent elsewhere in the English lexicon (except perhaps in other blends like *infotainment*; Lehrer, 2007), so such forms cannot be learned distributionally as regular combinative forms can. Nevertheless, language-users can reliably recognize the constituents *education* and *entertainment* in *edutainment* even without context.

Language-users might apply a general mechanism in which they attempt to match *tainment* to a list of possible source-words that contain this partial string (e.g., Lehrer, 1996). However, it is not totally clear—especially without context—how they would identify the source-word as *entertainment* and not *attainment* (which may be semantically more closely-related to *education*) or *containment*.

In this paper, we extend an amorphous model of morphological processing based on discriminative learning (Baayen, Milin, Filipovic Durdevic, Hendrix, & Marelli, 2011, and references within, summarized below) to explain how language-users recognize the source-words in blends like *edutainment*. Previous work has used this model to simulate lexical reading times in a way that accurately reflects a variety of known morphological processing phenomena (Baayen, 2010; Baayen et al., 2011; Baayen, Hendrix, & Ramscar, 2013). Here, we test human participants' ability to recover English blend constituents, and find that the model reliably predicts their success rate at guessing each source-word of a set of attested and novel blends. This empirical test provides evidence about the potential for this processing model to capture offline parsing intuitions in addition to online retrieval processes.

## Discriminative learning

In a discriminative learning model, individuals learn to acquire associations between CUES and OUTCOMES as a result of trials in which a cue co-occurs, or fails to co-occur, with a particular outcome. If a cue  $C_i$  occurs simultaneously with a particular outcome  $O$ , then a learner will increase their association weight  $V_i$  between the cue and outcome.

$$\Delta V_i = \alpha_i(\lambda - V_i) \quad (1)$$

In this equation,  $\alpha$  is a salience parameter, and  $\lambda$  is the theoretical maximum association weight that a learner can have between a cue and an outcome. If  $C_i$  occurs but the outcome  $O$  does not, the learner will decrease the association  $V_i$ .

$$\Delta V_i = \alpha_i(0 - V_i) \quad (2)$$

The size of each adjustment thus depends on the current value of  $V_i$ . As  $V_i$  gets closer to the theoretical maximum

association weight  $\lambda$ , each trial in which the cue and outcome co-occur will increase  $V_i$  by a smaller amount. As  $V_i$  gets larger, each trial in which the cue and outcome fail to co-occur will decrease  $V_i$  by a larger amount.

Crucially, when multiple cues occur together with an outcome, the learner does not treat them independently (Rescorla & Wagner, 1972). The adjustment to a cue  $C_i$  is dependent on the summed association weight between the outcome and the full set of cues that are present. Here,  $C_P$  represents the set of cues that are present simultaneously during the trial in which the outcome occurs, and  $\beta$  is a learning rate parameter. Equation 3 indicates the change in association between  $C_i$  and the outcome  $O$  when they occur together, and Equation 4 indicates the change in association when  $C_i$  occurs without outcome  $O$ .

$$\Delta V_i = \alpha_i \beta_1 (\lambda - \sum_{j \in C_P} V_j) \quad (3)$$

$$\Delta V_i = \alpha_i \beta_2 (0 - \sum_{j \in C_P} V_j) \quad (4)$$

For example, imagine a learning trial in which a rat sees a blue light and a red light—two cues—immediately before receiving an electric shock—the outcome. If the rat has no association between either light and the outcome, the trial will cause it to increase both association weights (red light  $\rightarrow$  shock, and blue light  $\rightarrow$  shock) by a large amount. In later trials, both cues will be considered to be good predictors of the shock. On the other hand, if the rat already strongly associates the red light and the shock, while it is seeing the blue light for the first time, the trial will cause only a small increase in association between each cue and the shock. In later trials, the blue light will not be considered as good a predictor of the shock, unless it continues to reliably occur with that outcome. The rat can learn to discriminate which cue is a good predictor of the shock (Ramscar, Yarlett, Dye, Denny, & Thorpe, 2010).

In the long term, an equilibrium point can be derived for each  $V$  if the following things are known:

1. probability  $P(O | C_i)$  of each outcome given each cue
2. probability  $P(C_j | C_i)$  of each cue given each other cue

In this naïve model, outcomes are considered to be independent of each other. For each outcome, the equilibrium weight  $V_i$  for each cue is found by solving the following system (Danks, 2003).

$$\begin{pmatrix} P(C_0 | C_0) & P(C_1 | C_0) & \dots & P(C_n | C_0) \\ P(C_0 | C_1) & P(C_1 | C_1) & \dots & P(C_n | C_1) \\ \vdots & \vdots & \ddots & \vdots \\ P(C_0 | C_n) & P(C_1 | C_n) & \dots & P(C_n | C_n) \end{pmatrix} \begin{pmatrix} V_0 \\ V_1 \\ \vdots \\ V_n \end{pmatrix} = \begin{pmatrix} P(O | C_0) \\ P(O | C_1) \\ \vdots \\ P(O | C_n) \end{pmatrix} \quad (5)$$

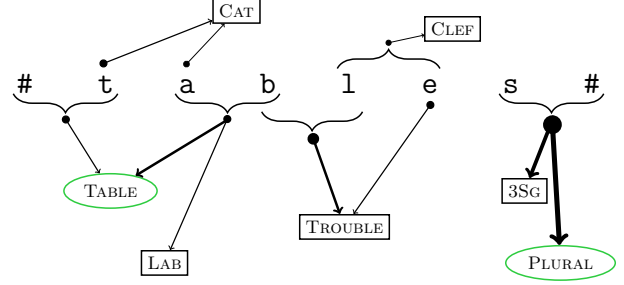


Figure 1: Association weights of different strengths (arrow thicknesses) between some of the unigram and bigram letter cues in the wordform *tables* and meaning outcomes (correct meanings in ellipses; other activated meanings in boxes)

## Morphological processing

In one linguistic version of the discriminative learning model, the CUES are taken to be short segment or letter strings, and the OUTCOMES are word meanings (Baayen et al., 2011; Ramscar et al., 2010). A language-user learns to associate each cue with each outcome—they acquire associations between symbols and their semantic knowledge of the world. If a segment or letter string frequently co-occurs with a particular lexical or grammatical meaning, but not elsewhere, the language-user will learn a strong association between that cue and that meaning.

For example, word-final *s* very often co-occurs in words with a PLURAL meaning, such as in *tables*, *books*, *chairs*, and many other nouns. Each time a language-user sees this co-occurrence, they more strongly associate word-final *s* with the meaning PLURAL. Word-final *s* also sometimes occurs without the PLURAL meaning, such as in *guess* or *mass*. When the language-user sees word-final *s* without the PLURAL meaning, they decrease their association between *s* and PLURAL. It is important to note that, in this model, the learner does not necessarily assign any privileged morphemic status to word-final *s*—they learn only that this cue is strongly associated with a particular meaning (Ramscar & Yarlett, 2007; Baayen et al., 2011, pp. 450–451).

In previous literature, cues are taken to be letter unigrams and bigrams. Word meanings typically include the lexeme name plus any inflectional information. For example, the meanings that occur with the word *tables* are TABLE and PLURAL. Corpus data can be used to estimate relatively how frequently each meaning occurs with each cue, and how frequently the cues occur with each other. These relative frequencies should be more or less stable in the long term, and so the system of equations in (5) can be used to derive the expected weights between the cues and outcomes. Figure 1 illustrates the relative strength of the equilibrium association weights between the cues in the form *tables* and a few possible outcomes.

When a reader encounters a word, the naïve discriminative learning model makes the following prediction. If the

summed total association weight between each cue in the word and the target outcome (the lexical meaning of the word) is high, it should be easier for the reader to retrieve the lexical representation of that word. In particular, the reader should read the word more quickly and they should have a more reliable judgment about the meaning of the word.

Previous literature has shown that the differential reading times predicted by a discriminative learning model account for a variety of known processing effects, such as inflectional regularity, frequency, and morphological family size (Baayen, 2010; Baayen et al., 2011, 2013). However, the model has yet to be evaluated on the reliability of morphological judgments—previous work has focused on response latencies. Further, it has yet to be used to predict either comprehension speed or reliability for the constituents of novel words.

### Modeling the recognition of blend constituents

Since the model does not attempt to decompose words into morphological constituents, it can straightforwardly explain how blend source-words are recovered. When a reader observes a string of letter cues, even if that string has never been seen before, those cues cause the activation of various meaning outcomes.

Therefore, the model makes a quantified prediction: language-users should be more likely to guess a source-word of a blend when there is a strong summed association weight between the orthographic cues in a blend and the meaning of that source-word. In other words, if the blend strongly activates one of its lexical source meanings, that lexical meaning should be easier to guess.

### Experiment 1: Attested blends

We asked English-speaking participants to guess the source-words of blends that are attested in English. We predicted that their aggregate success rate at guessing each blend source-word would be correlated with how strongly the blend activates the lexical meaning of that source-word.

**Procedure** 100 Mechanical Turk workers were paid \$0.25 for participation. Only participants with United States IP addresses and who certified that they were native speakers of American English were allowed to participate.

Each participant was presented with a random sample of 50 attested blends, with presentation order randomized for each participant. For each blend, participants were asked (1) to guess its two source-words and (2) to indicate whether or not they had seen the blend before in their prior experience.

**Stimuli** Attested blend stimuli were taken from the lists provided by Lehrer (2007) and Pound (1914). We attempted to address several possible confounds in the list of blends. If participants had seen a blend before, they might already know the source-words, either because they had been taught the sources, or because they had inferred them from context. Therefore, we excluded trials in which the participant indicated that they had seen the blend before. Further, we ex-

cluded blends that appeared in the Corpus of Contemporary American English (Davies, 2008) with a frequency greater than one per million.

Blends were also excluded if they contained a productive partial form, such as the *-holic* form in *workaholic*, *shopaholic*, *chocoholic*. Participants may have acquired new semantics for these forms (e.g., *-holic* is related to addiction, not to alcoholism per se), and language-users may decompose these forms as in normal derivation (Lehrer, 1998). A partial form was considered to be productive if it met either of the following criteria:

- it was categorized as a combining form in the appendix to Lehrer, 1998 (e.g., *-thon*, *-jacking*)
- it was named as a reusable form or possible bound morpheme in Lehrer, 2007 (e.g., *-umentary* in *mockumentary*)

Some further blends were not included in the stimuli set because it seemed unlikely that any participant would successfully guess the source-words. In particular, blends were excluded if either source-word had a frequency of less than one per million in COCA, or if they involved words or concepts not in contemporary use (e.g., *terrible* > *torrid* + *horrible*). Finally, blends were excluded if either source-word did not appear in the English CELEX database used to train the discriminative learning model.

Attested blends were also excluded if they were proper names (*Craisins*); or were not composed of nouns, adjectives, or verbs (*thon* > *that* + *one*); or were composed of more than two source-words (*skafrocuban* > *ska* + *Afro* + *Cuban*).

This left a final list of 89 attested blends used as stimuli in Experiment 1.

**Model** Using the *ndl* package for R (Arppe, Milin, & Hendrix, 2012), a naïve discriminative reading (NDR) model was trained on the wordforms and frequency data in the English CELEX database (Baayen, Piepenbrock, & Gulikers, 1996). The meanings (outcomes) associated with each wordform included the lexeme name in addition to the inflectional meanings provided in the morphological wordform annotations. For example, the meanings associated with the form *geese* were GOOSE and PLURAL, and the meanings associated with the form *driven* were DRIVE, PARTICIPLE, and PAST. As in previous work, letter unigrams and bigrams were used as cues to meaning. For example, the cues that appear in the form *geese* are: g, e, s, #g, ge, ee, es, se, e#.

**Activation strength** The trained NDR model was used to calculate the total activation strength between the cues in each blend and the meanings of its source-words. The strength with which a blend activated a target meaning was considered to be the sum of the association weights between that lexical meaning and each cue in the blend. We predicted that the total activation of a source-word meaning like ENTERTAINMENT by the cues in *edutainment* should be correlated with how easily participants are able to guess that *entertainment* is one of the source-words of *edutainment*.

**Control variables** Based on previous literature, we expected that a number of other factors would influence how easily participants would be able to guess a blend source-word (Lehrer, 1996). Control variables included:

- the number and percentage of letters from the target source-word that were retained in the blend
- the source-word’s frequency in CELEX
- the number of word lemmas in CELEX that the orthographic partial form could possibly have been taken from (e.g., how many words could *tainment* possibly be from)
- the ratio of the frequency of the correct source-word to the summed frequency of all other possible source-words
- the average probability of letter trigrams in the source-word, according to an orthographic model trained on the Brown corpus
- the source-word’s *minimum* orthographic probability using the same metric (following e.g., Hay & Baayen, 2003, who argue that low phonotactic probability at a boundary facilitates the perception of morphological complexity)
- whether the participant correctly guessed the other source-word in the blend
- whether the source-word occurred second in the blend

**Results** The lme4 package for R (Bates, Maechler, & Dai, 2008) was used to fit a mixed-effects logistic regression predicting participants’ success rate at identifying each source-word on the basis of the blend’s NDR activation of the source-word meaning and the control variables listed above. A source-word identification was marked as correct if a participant guessed either the target word or an inflected form, but was marked as incorrect for other wordforms, including derived forms containing the target word. Responses were excluded if the participant indicated that they had seen the blend before, if the response was left blank, or if the response suggested that the participant misunderstood the task (e.g., writing *animal + king* as the source-words of *zebrule*). The final analysis included 8,008 source-word guesses (72% correct).

The model also included per-subject and per-source-word random intercepts and slopes where they were justified by the design, including per-subject NDR activation slopes and the maximal structure that allowed the model to converge. The results of the logistic regression are presented in Figure 2(a).

Crucially, NDR activation was found to be reliably predictive of participants’ ability to guess blend sources ( $\beta = 0.53$ ,  $z = 2.6$ ,  $p < 0.01$ ). Four control factors were also significant. The more material from the target source-word retained in the blend (measured in both raw number of letters and percentage of letters), the more likely participants were to correctly guess the source-word ( $\beta = 0.49$ ,  $z = 2.6$ ,  $p < 0.01$ ;  $\beta = 0.97$ ,  $z = 5.7$ ,  $p < 0.0001$ , respectively). Additionally, participants were more likely to guess a source-word

correctly if they correctly guessed the blend’s *other* source-word ( $\beta = 1.65$ ,  $z = 8.5$ ,  $p < 0.0001$ ). Finally, source-words with higher CELEX word frequency were more likely to be guessed correctly ( $\beta = 0.79$ ,  $z = 2.3$ ,  $p < 0.03$ ).

## Experiment 2: Novel blends

One possible concern with using attested blends is that many of them survived as lexical items for a relatively long time. It may be the case that these blends are exceptional in some way. For example, they might remain in use because they are unusually easy to parse or understand, which would be a potential confound for the results in Experiment 1. Therefore, we conducted a second experiment using novel blends that were constructed specifically for the experimental task.

**Stimuli** To serve as the source-words for novel blends, 20 pairs of co-hyponyms were selected from WordNet (Fellbaum, 1998). The co-hyponymy relationship has been argued to be one of the most common semantic relationships between the two source-words in a blend (Gries, 2012). In each word, we marked the boundaries between orthographic syllables, and between the onset and rime, as possible split points. This was done to improve the phonological well-formedness of the resulting blends; the onset-rime boundary has been shown to be a common split point for blends, at least in English (Kelly, 1998; Gries, 2004).

To construct blend stimuli for each pair, one of the split points was randomly selected in each source-word. For example, in the pair *insult* and *sting*, we might select the boundaries *in—sult* (between the two syllables) and *—sting* (word-initial). The material preceding the split point in one of the words was concatenated with the material following the split point in the other word, based on a random coin toss. This procedure only has the possibility of creating linear blends, and excludes possibilities like *chortle* > *chuckle* + *snort*, in which two disjoint parts of the first source are separated by a word-medial piece of the second source. Novel blends were also constrained by the requirements to contain at least one unique letter from each source-word, to contain at least one orthographic vowel, to be shorter than the concatenation of both full sources (i.e., no full compounds), and to be longer than the shorter source-word.

Using this method, four possible blends were generated for each source-word pair, resulting in 80 total blends.

Table 1: Sample novel blends.

| source pair                          | novel blends                                                               |
|--------------------------------------|----------------------------------------------------------------------------|
| { <i>insult</i> , <i>sting</i> }     | <i>insting</i> , <i>stingsult</i> , <i>stingult</i> , <i>stult</i>         |
| { <i>sofa</i> , <i>stool</i> }       | <i>sofool</i> , <i>sool</i> , <i>sostool</i> , <i>stoolfa</i>              |
| { <i>diagram</i> , <i>scribble</i> } | <i>diagribble</i> , <i>scribbagram</i> , <i>scribbam</i> , <i>scrigram</i> |

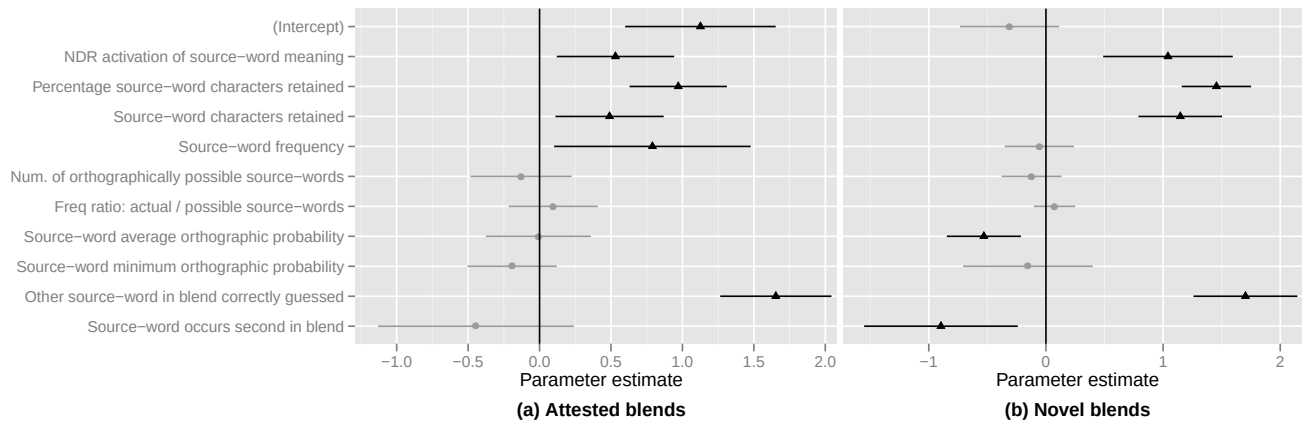


Figure 2: Results of logistic regression predicting human success rate in guessing source-words of (a) attested blends and (b) novel blends. All continuous variables were centered and standardized. Points are parameter estimates  $\beta$ ; bars reflect two standard errors. Significant factors appear in black (reported  $p$  values are based on the Wald  $z$  statistic). For both experiments, backward model selection was also performed to remove non-significant predictors, but the remaining significant effects were found to be qualitatively the same as in the full models. For Experiment 1, condition number  $\kappa = 4.8$ ; for Experiment 2,  $\kappa = 3.8$ , which both indicate low collinearity.

**Procedure** The procedure was identical to that of Experiment 1, except that participants were not asked whether they had seen the blend before. 100 new participants were recruited. For each source-word pair, each participant saw exactly one of the four novel blends, chosen at random.

**Results** We fit a mixed-effects logistic regression, using the same procedure as in Experiment 1, to predict participants' success rate at identifying each source-word on the basis of the blend's NDR activation of the source-word meaning in addition to the same control variables. Random per-blend intercepts and slopes were also added to the model structure, since there were multiple blends for each source-word. With the exclusions described above, 3,736 source-word guesses (50% correct) were included in the final analysis.<sup>1</sup> The results of this model are presented in Figure 2(b).

NDR activation was again found to be reliably predictive of participants' ability to guess blend sources ( $\beta = 1.04$ ,  $z = 3.8$ ,  $p < 0.001$ ). All of the significant control factors for attested blends were significant for novels, except source-word frequency (letters retained:  $\beta = 1.15$ ,  $z = 6.4$ ,  $p < 0.0001$ ; percentage of letters retained:  $\beta = 1.46$ ,  $z = 9.8$ ,  $p < 0.0001$ ; other source-word guessed correctly:  $\beta = 1.70$ ,  $z = 7.7$ ,  $p < 0.0001$ ). The difference in the frequency effect may be due to the different stimuli: there were only 40 source-words for the novel blends, and these were hand-selected and mostly of medium frequency.

Two additional factors were significant for novel blends. First, the higher the average orthographic probability of the

source-word, the less likely it was to be guessed ( $\beta = -0.53$ ,  $z = -3.4$ ,  $p < 0.001$ ). In other words, the more orthographically unusual the source-word, the more likely participants were to recover it. Second, source-words in second position in the blend were less likely to be recovered than those in first position ( $\beta = -0.90$ ,  $z = -2.7$ ,  $p < 0.01$ ). This may suggest that word-initial letters are slightly more salient to readers than medial or final letters.

## Discussion

The amorphous letter-to-meaning associations of the naïve discriminative reading model were found to be good predictors of human readers' success at recovering the constituents of both attested and novel blends. If a blend contained orthographic cues that strongly activated the lexical meaning of one of its source-words, readers were better able to recognize that source-word in the blend form. This effect was independent of source-word frequency, the length of the partial form, and a number of other control variables.

## Discriminative learning and constituent recognition

The result extends previous literature, which has argued that an NDR-based model can account for morphological phenomena associated with reading times. In particular, the current study supports a distinct prediction of the discriminative learning account: meaning activations are correlated with the reliability and success with which readers recover morphological constituents offline, in addition to their speed at doing so. Additionally, previous work has looked primarily at word-forms and constructions that are already known to the learner. This study extends this work by showing that the model can also account for comprehension effects in forms that a reader

<sup>1</sup>Responses to one source-word pair (*strength*, *advantage*) were also excluded because these responses unusually inflated the variance of the random effects estimates; however, all results were qualitatively the same with these responses included.

has not previously been exposed to. This provides evidence that discriminative learning captures the effects of processing mechanisms, rather than stored representations of existing words.

This implementation of the NDR model describes the recovery of the component meanings of word-forms and constructions in isolation, but does not purport to provide an account of inference of whole-form meaning. For example, the blend *dogbrella* might mean “an umbrella for dogs” or “an umbrella with pictures of dogs.” Pragmatic and linguistic context as well as prior probability distributions over semantic relationships (Pollatsek, Drieghe, Stockall, & de Almeida, 2010) provide a basis for future extension of the model.

### Blend processing

The NDR model is designed to capture form-to-meaning relationships without requiring formal decomposition. In particular, it allows the extraction of multiple (lexical or grammatical) meanings from non-decomposable forms with multiple constituents. Blends are an excellent test case for such a model, because they conspicuously have multiple constituents while at the same time they are not subject to any kind of regular decompositional analysis. A purely-compositional account such as that of Marantz (2013) is challenged to explain how a conjunction of multiple partial lexical wordforms can be reliably parsed into its original constituents. On our account, in contrast, blend processing is modeled as prediction of meanings based on many small cues, namely all (unigrams and bigrams of) letters in the blend.

The discriminative learning model may help explain how language-users recover *entertainment* from the partial form *tainment* instead of plausible alternatives like *attainment* or *containment*. The cues in *edutainment* are collectively better associated with *entertainment* than the alternatives, which leads to a stronger activation of *entertainment* and thus a greater likelihood that a reader will select the correct form.

Our account additionally makes tractable predictions about the structure and function of blends in natural language. One communicative constraint on the formation of blends might be that the meanings of individual source-words must be recoverable by comprehenders (Lehrer, 1996; Gries, 2004). If this is true, we predict that cues (here, letter unigrams and bigrams) that most strongly activate the intended meanings are most likely to become part of the blend, and that language-users would judge blends with these cues to be better than blends with cues that less strongly activate these meanings. This is the subject of ongoing research.

### Acknowledgments

We are grateful to Roger Levy, Farrell Ackerman, Robert Malouf, Ryan Lopic, David Barner, Olivier Bonami, the UCSD CPL Lab, the morphology group at UCSD, and the audience at BLS 40 for insightful discussion and suggestions. We also thank Bill Presant for assistance with data annotation. Any errors or omissions are ours. This work was supported by NSF Graduate Research Fellowships to both authors.

### References

- Arppe, A., Milin, P., & Hendrix, P. (2012). ndl: Naive discriminative learning [Computer program]. (R package version 0.1.6)
- Baayen, R. H. (2010). Demythologizing the word frequency effect: A discriminative learning perspective. *The Mental Lexicon*, 5(3), 436–461.
- Baayen, R. H., Hendrix, P., & Ramscar, M. (2013). Sidestepping the combinatorial explosion: An explanation of n-gram frequency effects based on naive discriminative learning. *Language and Speech*, 56(3), 329–347.
- Baayen, R. H., Milin, P., Filipovic Durdevic, D., Hendrix, P., & Marelli, M. (2011). An amorphous model for morphological processing in visual comprehension based on naive discriminative learning. *Psychological Review*, 118, 438–482.
- Baayen, R. H., Piepenbrock, R., & Gulikers, L. (1996). *CELEX-2* (Tech. Rep.). Philadelphia: Linguistic Data Consortium.
- Bates, D., Maechler, M., & Dai, B. (2008). lme4: Linear mixed-effects models using s4 classes [Computer program]. (R package version 0.999375-28)
- Danks, D. (2003). Equilibria of the Rescorla-Wagner model. *Journal of Mathematical Psychology*, 47, 109–121.
- Davies, M. (2008). *The Corpus of Contemporary American English: 450 million words, 1990–present*. (<http://corpus.byu.edu/coca/>)
- Fellbaum, C. (1998). *WordNet: An electronic lexical database*. Cambridge, MA.
- Finley, S., & Newport, E. L. (2011). Morpheme segmentation in school-aged children. *University of Rochester Working Papers in the Language Sciences*.
- Finley, S., & Wiemers, E. (2013). Rapid learning of morphological paradigms. In *Proceedings of the 34th Annual Conference of the Cognitive Science Society*. Austin: Cognitive Science Society.
- Gries, S. T. (2004). Shouldn't it be breakfunch? A quantitative analysis of blend structure in English. *Linguistics*, 639–668.
- Gries, S. T. (2012). Quantitative corpus data on blend formation: psycho- and cognitive-linguistic perspectives. In V. Renner, F. Maniez, & P. Arnaud (Eds.), *Cross-disciplinary perspectives on lexical blending* (pp. 145–167). New York: Mouton de Gruyter.
- Hay, J., & Baayen, H. (2003). Phonotactics, parsing and productivity. *Italian Journal of Linguistics*, 1, 99–130.
- Kelly, M. H. (1998). To ‘brunch’ or to ‘branch’: Some aspects of blend structure. *Linguistics*, 36, 579–590.
- Lehrer, A. (1996). Identifying and interpreting blends: an experimental approach. *Cognitive Linguistics*, 7(4), 359–390.
- Lehrer, A. (1998). Scapes, holics, and thons: The semantics of English combining forms. *American Speech*, 73(1), 3–28.
- Lehrer, A. (2007). Blendalicious. In *Lexical creativity, texts and contexts* (pp. 115–133). Amsterdam: John Benjamins.
- Marantz, A. (2013). No escape from morphemes in morphological processing. *Language and Cognitive Processes*, 28(7), 905–916.
- Pollatsek, A., Drieghe, D., Stockall, L., & de Almeida, R. (2010). The interpretation of ambiguous trimorphemic words in sentence context. *Psychonomic Bulletin & Review*, 17(1), 88–94.
- Pound, L. (1914). *Blends: their relation to english word formation*. Heidelberg: Winter.
- Ramscar, M., & Yarlett, D. (2007). Linguistic self-correction in the absence of feedback: a new approach to the logical problem of language acquisition. *Cognitive Science*, 31(6), 927–60.
- Ramscar, M., Yarlett, D., Dye, M., Denny, K., & Thorpe, K. (2010). The effects of feature-label-order and their implications for symbolic learning. *Cognitive Science*, 34(6), 909–57.
- Rescorla, R. A., & Wagner, A. R. (1972). A theory of pavlovian conditioning: Variations in the effectiveness of reinforcement and nonreinforcement. In A. H. Black & W. F. Prokasy (Eds.), *Classical Conditioning II: Current Research and Theory* (pp. 64–99). New York: Appleton-Century-Crofts.
- Stockall, L., & Marantz, A. (2006). A single route, full decomposition model of morphological complexity: MEG evidence. *The Mental Lexicon*, 1(1).
- Taft, M. (1981). Prefix stripping revisited. *Journal of Verbal Learning and Verbal Behavior*, 20(3), 289–297.
- Taft, M., & Forster, K. I. (1975). Lexical storage and retrieval of prefixed words. *Journal of Verbal Learning and Verbal Behavior*, 14, 638–647.