

Accounting for underlying forms in HPSG phonology

Filip Skwarski

University of Warsaw

Proceedings of the 16th International Conference on
Head-Driven Phrase Structure Grammar

Georg-August-Universität Göttingen, Germany

Stefan Müller (Editor)

2009

Stanford, CA: CSLI Publications

pages 318–337

Skwarski, Filip. 2009. Accounting for underlying forms in HPSG phonology.
In Stefan Müller (ed.), *Proceedings of the 16th International Conference on Head-Driven Phrase Structure Grammar, Georg-August-Universität Göttingen, Germany*, 318–337. Stanford, CA: CSLI Publications. DOI: 10.21248/hpsg.2009.16. 
CC-BY 4.0

Abstract

The paper aims to present approach to HPSG phonology which would account for underlying forms of phonemes. It shows some of the issues arising in monostratal analyses of phonology, and proposes a solution based on a notion of underlying representations. The approach presented, partly inspired by Optimality Theory, resolves cases of neutralisation and opacity by formulating constraints which either restrict the surface representation or relate it to the underlying form.

1 Why are underlying forms desirable in HPSG phonology?

Since most work in HPSG is focused on syntax or morphology, standard representations of phonology are reduced to supplying a word with a feature of type *phon*, a string of symbols equivalent to the orthographical spelling or pronunciation of the word (and its particular variants). Such strings are afterwards combined in higher-level objects to form phrases. While such a simplified approach is sufficient for solving problems based around syntax, morphology, and semantics, the HPSG does show a lot of potential for expanding the phonology within its framework.

One notable attempt to do so was undertaken by Bird (1995), who introduced constraint-based phonology into HPSG. The framework outlined in *Constraint Based Phonology: A Computational Approach* is essentially a monostratal system, where well-formedness of a particular word or phrase is decided by phonological and morphological constraints. This system allowed for linking phonology and morphology, and resolved issues by ruling out ill-formed segments and word structures.

The framework proposed was based around the principles of COMPOSITIONALITY and a requirement that a framework be MONOSTRATAL. The latter meant, in simplified terms, that any phonological representation has only one level, corresponding to forms actually appearing in the surface representations, and no abstract representation is stored:

An even stronger constraint than those mentioned above is the requirement that a linguistic framework be MONOSTRATAL. This means that there is only a single level of linguistic description; descriptions pertain to occurring surface forms and not to artificially constructed abstract representations. As we shall see in section 1.5.1, the requirement that a linguistic framework be monostratal is equivalent to

the true generalisation condition from Natural Generative Phonology.
(Bird 1995, 1.4.5, p. 34)

Although such an approach would seem to be desirable in a computational framework, the phonological phenomena in various languages cannot be adequately described without a further reference to an underlying representation of a phoneme (Shoun 2005, 4.4). The cases pointed out in Shoun (2005) include eg. neutralisation phenomena in Bengali, there is no proposal as for the actual implementation of the underlying forms in HPSG phonology, however - which is the aim of this paper.

Evidence for usefulness of underlying representations can be seen in consonant alternations and voicing processes in languages where those phenomena are complicated, even though Bird (1995) seems to disregard events such as final devoicing as purely phonetic processes which need not be described with binary features, based on Port and Crawford (1989):

The data show that speakers can control the degree of neutralisation depending on pragmatics and that information about the underlying contrast is distributed over much of the word. The data support a scalar valued neutralisation effect in the German voicing rule, and clearly refute a rule using a binary voicing feature. (Port and Crawford 1989, 257f)

Assuming such a position avoids the problem entirely by postulating that no instances of homophony due to devoicing exist, or in general that many alternation-related phenomena can be simplified as phonetic processes, while a substantial amount of evidence points out to the contrary.

In Polish (my native language), for example, the phoneme /g/ exhibits the following alternations:

(1)	księga	a tome (nom.sg.)	[kɕɛŋga]	[g]
	ksiąg	of tomes (gen.pl)	[kɕɔŋk]	[k]
	księdze	to a tome (dat.sg.)	[kɕɛndzɛ]	[dz]
	książka	a book (nom.sg.)	[kɕɔ̃ʦka]	[ʃ]
	książek	of books (gen.pl.)	[kɕɔ̃ʦɛk]	[ʒ]

Although these alternations result from historical palatalisation and voice assimilation processes, all of them are fully productive in modern Polish, in specific morphology-related cases, like noun declension patterns.

Likewise, in Polish - unlike eg. German - the process traditionally called "final obstruent devoicing" is intertwined with a process of "voice assimilation". Voiced obstruents are devoiced word-finally and before voiceless obstruents, while voiceless obstruents become voiced before voiced obstruents, including across word boundaries (Rubach 1982, 4.2, 4.3). As a result, /d/ and /t/ can both surface as [t] and [d] accordingly, phonetically identical with the "default" form of their opposite-voiced counterpart. Before sonorants (except, in most cases, across word boundaries), obstruents retain their "underlying" voice values, and so, in a traditional monosyllabic framework, we would have no way of arriving at this basic form if we simply described sonorants as either alternations of their surface representations or underspecifications (as suggested by Bird (1995), 1.5).

(2a)	kod	code	[kɔt]
	kody	codes	[kɔdɨ]
	kod dostępu	access code	[kɔd dɔstɛmpu]
	kod miasta	city code	[kɔt mʲasta]
	kod pocztowy	postal code	[kɔt pɔtʃtɔvɨ]
(2b)	kot	cat	[kɔt]
	koty	cats	[kɔtɨ]
	kot domowy	house cat	[kɔd dɔmɔvɨ]
	kot mały	small cat	[kɔt mawɨ]
	kod perski	Persian cat	[kɔt pɛrski]

The above data demonstrates that obstruents in Polish can behave in three ways depending on context:

1. assimilate their voice to that of the following segment (before other obstruents, including across word boundaries),
2. retain their "underlying" voice feature (before sonorants, except word-finally),
3. become voiceless regardless of their "underlying" voice feature (word-finally before a pause or before sonorants).

This example (not unlike the neutralisation example in Bengali, in Shoun 2005) will be used as a basis for representing the possibilities of accounting for underlying forms in HPSG phonology, in a further section.

1.1 Underspecification and Surface Constraints

Before presenting an approach utilising the notion of underlying representation to resolve these issues, it is worth looking at some of the views on alternations in HPSG phonology presented so far. One of the possibilities in line with Bird's orig-

inal idea would be to use purely surface-relevant constraints, providing separate structures for various levels of the sentence if necessary (words, utterances, syllables, etc.). Such a solution is adopted by Bird and Klein (1994) and suggested by Höhle (1999).

There are two possibilities of expressing the phenomenon of Polish final devoicing within such a framework:

1. Restrict the word structure in such a way that no voiced word-final obstruents are allowed, and the phrase structure in such a way that all obstruent segments must agree in voicing.
2. Restrict the word structure in such a way that no voiced word-final obstruents are allowed before a sonorant, and the phrase structure in such a way that all obstruent segments must agree with voicing.

Of these, solution 1. leads to an obvious conflict whereby in a phrase "kod dostępu" the phrase demands a voiced [d] while the word demands a voiceless [t], and therefore no proper form can be generated. Solution 2. leads to an underspecification, where in the cases of "kod", "kot", "kod pocztowy" or "kod dostępu", the voicing value is correctly predicted, but in "kody" or "koty" (word-medially), it is not determined at all (it is [d] v [t]), and we are in fact left with no means to predict it. We simply cannot "consult" it with anything.

Höhle (1999) attempts to tackle Russian obstruent voicing rules, not very different from the Polish ones, and in his approach seems to allow for different phoneme surface representations arising on different levels (Höhle 1999, fig. 7). While it is possible to differentiate between the representations of a particular phoneme present on the word level and the phrase level (by simply not making them identical in HPSG sense, and by arriving at the two by separate means), this leads to problems with coordinating the entire structure - see the notes on principles adopted for this framework in section 2., "Representing the Representations".

1.2 Morphology and Stem Spaces

Apart from Bird's and Höhle's proposals regarding HPSG phonology, an attempt to tackle phonological alternations (including irregular patterns displayed by some forms) is demonstrated in *Deriving Inflectional Irregularity*, Bonami and Boyé (2006). Here, a notion of stem space is introduced: declension patterns are based around a number of stem spaces, accounting for all stem alternations in inflection. For example, instead of (as assumed in transformational phonology) deleting endings of French feminine adjectives to produce masculine ones, the two are assumed to have different basic stem spaces (Bonami and Boyé 2006, 2.2).

But while French is (relatively) simple in terms of phonological processes, adopting such a framework in Polish would be complicated for a number of reasons:

1. In Polish, the voicing phenomena - as demonstrated above - are not only affected by the declension pattern, but also by the context of the following and preceding words. Eg. the lexical item "kod" cannot be simply a part of a class of nouns where [d] occurs in the stem before affixes and [t] in the nominative, unless we further account for the fact that the [t] in the nominative may alternate with [d] anyway.

Moreover, entire clusters may change their voice features:

(3a)	mózg	a brain	[musk]
	mózg był	a brain was	[muzg bɨw]
(3b)	zjadł	he ate	[zjatɯ]
	zjadł go	he ate him	[zjadw gɔ]

2. Polish is further complicated by other phonological phenomena accounting for further alternations:

(4a)	kocha	loves	[kɔxa]
	kochają	(they) love	[kɔxa jɔ̃w]
	kochając	loving	[kɔxa jɔnts]
(4b)	robi	does	[rɔbʲi]
	rób	do (imp.)	[rup]
(4c)	zjedli	they (masc.) ate	[zjɛdli]
	zjadliwe	edible (neut.)	[zjadlive]
	zjadł	he ate	[zjatɯ]
	zjedzony	eaten (masc.)	[zjɛdzɔni]
(4d)	Paryż	Paris	[parɨʃ]
	paryski	Parisian (masc.)	[pariski]

The above examples demonstrate some of the processes: alternation between nasal vowels followed by a glide and oral vowels followed by a nasal consonant in (4a), alternation between [u] and [ɔ] in (4b), alternation between [l] and [w], [ɛ] and [a], as well as [d] and [dz] in (4c), and disappearance of fricatives in (4d). In addition, all obstruents (and extra-syllabic approximants) are also affected by voicing rules, best exemplified in (4c).

In the more extreme cases like (4c), virtually any segment found in a word can alternate with something else, leading to some situations where representing alternations exhibited by individual phonemes is actually easier than representing alternations of all possible stem forms.

3. Approaches based on morphological tools, unlike those based on global phonological constraints, essentially say nothing about the permissible structures of the words themselves, as long as they are assigned into productive patterns. In an optimal system, aside from handling alternations, it would be desirable to predict which words are well-formed according to the phonotactics of a given language, especially since, as demonstrated before, global constraints are useful in ruling out erroneous forms in situations where the choice of a proper form of a stem, affix, etc., is based purely on phonological background. On the other hand, introducing separate information about phonotactics would in many cases overlap with what is already handled by morphology.

The above remarks should demonstrate the major problems arising when using morphology-based tools and noun classes for adequately representing actual pronunciation of spoken words. While such an approach could be expanded, it would have to become mind-bogglingly complex for some languages, while invoking underlying representation removes the need for merging morphological, phonological and phonostylistic phenomena into one monster of a framework - every process can be dealt with separately by operating at the level at which it occurs and on the phonemes or features it is related to. No distinction between phonotactics and inflection becomes necessary, even though more advanced morphological issues, like the examples back in (1), can be addressed by invoking both morphological classes and the underlying representations.

2 Representing the Representations

This section is concerned with establishing the structural side of the framework which would involve underlying features. A well-functional framework should achieve the following aims:

- (a) Allow formulated rules to operate at various levels of the structure (stem, word, syllable, utterance, etc.)
- (b) Accurately provide just one surface form for any phoneme in the complete utterance.
- (c) Append lower-level representations in higher-level representations (words into phrases, syllables into feet, etc.)
- (d) Allow for interactions between the underlying representation and the surface

representation in cases where the underlying representation is directly relevant to the surfacing form.

Principle (a) is dictated by the observation that certain phonological phenomena are limited within syntactic context, such as word boundaries or phrase boundaries, and the constraints have to be formulated in a way accounting for this (Bird 1994, 2.2).

Principles (b) and (c) are related: because of the observation mentioned in (a), various constraints operating solely on one level of the structure (word, phrase, etc.) would predict different criteria of well-formedness. For example, a constraint demanding that the word-final segment be voiceless would apply to the PHON structure of a *word* object, but not to the PHON structure of a *phrase* object. Similarly, constraints operating across word boundaries would not say anything about the PHON structure of a word object.

As a result, for a situation like the exemplary interaction between Polish final devoicing and voice assimilation processes (2a and 2b), we are left with a choice of either predicting different phonological structures for different levels of syntactic and morphological representation, or postulating that all surface representations at all levels have to be the same. Höhle (1999) appears to use (presumably for simplification) the first case scenario, and in his representations, different phoneme sorts (used for contrastive voicing) appear at the level of the word and at the level of an utterance. Applying this to our Polish devoicing example would yield a situation in which the phrase "kod dostępu" would have a PHON listing [kɔd dɔstɛmpu], but its first daughter element would still display a structure ending in a voiceless obstruent: [kɔt].

Such a solution makes it possible to account for predictions made at different levels, but causes problems with principle (c), that is, it requires a separate system for appending daughter elements together (since we cannot simply append [kɔt] and [dɔstɛmpu] to get [kɔd dɔstɛmpu]). Again, introducing underlying representation seems to be an advantage here, as it does not require clearly defined and sorted phonemes (which would be superfluous, Shoun 2005, 4.2), but allows forms to combine precisely because higher level structures are appended based on the underlying structure of their elements, while the surface structure may be separately predicted.

2.1 Summary of Proposals

To summarise - a system I propose is a system where the underlying and the surface forms are stored separately, where the higher level lists are appended separately

(from underlying and surface lists of daughter elements), and where the underlying and surface forms can interact through formulated rules and constraints. The following section shall encompass all the major technical details of such an approach.

The approach presented here uses two levels of phonological representation, but is otherwise non-transformational and does not rely on rule ordering. The notion of underlying representation goes way back to transformational phonology, but can be found also as a solid basis in more recent phonological theories, most notably Optimality Theory. Applied in HPSG, it would not produce the surface forms through ordered rules or evaluating a number of universal constraints, but by allowing constraints that relate the surface representation to the underlying one. Constraints could be formed involving either of the representations (underlying and surface), but the surface representation could be restricted to depend on the underlying representation in cases where it cannot be arrived at purely through surface level constraints.

The way the relationship between the underlying representation and the surface representation operates essentially resembles the core ideas of Optimality Theory (Prince and Smolensky 1993/2004, 1.2), where surface representations are selected depending on the criteria of surface well-formedness (markedness) and closeness to the underlying representations (faithfulness). However, in this HPSG-based approach, the relationship is unambiguous and rather than relying on an algorithm selecting the most favourable form according to a universal constraint hierarchy (as is the case in OT), the surface forms are predicted based on language-specific, global constraints.

The notion of the underlying representation adopted here adheres to that in Shoun (2005, 4.2), ie. there is no clearly defined "list" of underlying phonemes for a language. Segments which show no productive alternations within a stem can be represented as identical to the surface form, disregarding baroque historical recreations. Finally, some of the core structural and technical sides of this approach are based on Bird's original proposals (Bird 1995).

2.2 Organising Segments

Although phonology in HPSG is traditionally handled through lists, the solution I propose is to use a new type of object, which I term here *segs* (for "segments"). This seemingly bizarre decision is dictated by the aforementioned principles: in order to coordinate the PHON values of utterances, phrases, words, etc., and at the same time allow constraints to operate at different levels, *segs* can be divided into subtypes, ie. *utterance-segs*, *phrase-segs*, *word-segs*, etc. Furthermore, due to the

implementation of underlying representation, *segs* contains list features (UR-LIST and SR-LIST) for coordinating daughter elements, similar to DTR-LIST used by HPSG phrases.

The structure of *segs* would look like the following:

$$(5) \left[\begin{array}{ll} \text{segs} & \\ \text{SR-LIST} & \text{list} \\ \text{UR-LIST} & \text{list} \\ \text{FIRST} & \text{ph-str} \\ \text{REST} & \text{segs} \vee \text{e-list} \end{array} \right]$$

(*ph-str* here stands for "phonetic structure", and corresponds to the structure used to express the relationship between the UR and the SR, ie. one-to-one, one-to-many, or one-to-none)

While the FIRST feature of *segs* always has to be a phonetic structure, the REST can either be another *segs* or an empty list. Such a selection of REST value is not the only option: my original concept was to allow either *segs* or *ph-str*, in a manner resembling how phrases and words (or morphemes) are handled in HPSG syntax. However, such an approach requires us to either formulate numerous constraints twice, or introduce a phonological equivalent of Head Feature Principle. It is thus easier to settle down for ending all final *segs* in an empty list. While this adheres to the conventional way of handling list-like objects, it may cause its own problems with implementation by demanding an object which, in traditional HPSG ontology, belongs to an entirely different class than *segs*. One more possibility would be to replace *e-list* with a new, feature-empty subtype of *ph-str*, but for simplicity's sake I will just use the familiar *e-list* throughout the paper for *segs* and other list-like structures.

(in reality, the detailed structuring of *segs* is not as crucial as it seems, because most phonological structures can be introduced into the lexicon by specifying just UR-LIST, as shall be seen further on)

Last but not least, note that the features FIRST and REST are named after lists. In reality, HPSG ontology would demand these to be named distinctly in order to differentiate *segs* from regular lists: S-FIRST and S-REST are one of the possibilities, but I will use FIRST and REST throughout, again, for the sake of simplicity.

As an example of subtypes, in this model, the PHON feature of the word object will be an object of type *word-segs*, whose non-empty daughter elements will also have to be *word-segs*. But if the *word* object and eg. another *word* object get appended into a higher level *phrase* object, the PHON value of that phrase will be composed of *phrase-segs* objects.

$segs \rightarrow word-segs \vee phrase-segs \vee utterance-segs \dots etc.$

$$(6a) \text{ word} \rightarrow \begin{bmatrix} \text{word} \\ \text{PHON} \mid \text{SEG-LIST } \text{word-segs} \end{bmatrix}$$

$$(6b) \begin{bmatrix} \text{word-segs} \\ \text{REST } segs \end{bmatrix} \rightarrow \begin{bmatrix} \text{word-segs} \\ \text{REST } \text{word-segs} \end{bmatrix}$$

The structures for every level of the sentence will be, thus, different, but the constraints operating on UR-LIST and SR-LIST will demand that the actual phonetic representations in the daughter elements (stored in eg. *word-segs*: UR-LIST) be carried over and appended into mother elements (eg. *phrase-segs*: UR-LIST).

The correspondences between the underlying and the surface representations are handled through an object of type *ph-str* ("phonetic structure"), of which I propose three subtypes:

$$(8) \text{ ph-str} \rightarrow \begin{bmatrix} \text{simple} \\ \text{UR } rep \\ \text{SR } rep \end{bmatrix} \vee \begin{bmatrix} \text{complex} \\ \text{UR } rep \\ \text{SR} \begin{bmatrix} \text{complex-rep} \\ \text{SR-LIST } list \\ \text{FIRST } rep \\ \text{REST } \text{complex-rep} \vee e\text{-list} \end{bmatrix} \end{bmatrix} \vee \begin{bmatrix} \text{empty} \\ \text{UR } rep \end{bmatrix}$$

(*rep* stands for "representation": the actual phonetic description of features)

The *simple* object corresponds to the casual scenario where one underlying form corresponds to one uttered segment ("phone"). The *complex* object accounts for epenthesis, a case where one underlying phoneme corresponds to a more complex phonetic structure of two or more segments. Finally, the *empty* object accounts for deletion, ie. a situation the underlying segment is not visible on the surface at all. The two latter objects will be seen in action in section 2.2. on opacity in Turkish.

The formulation rules of UR-LIST and SR-LIST are expressed through the following constraints (for all three subtypes respectively):

$$(7a) \begin{bmatrix} \text{segs} \\ \text{FIRST } \text{simple} \\ \text{REST } segs \end{bmatrix} \rightarrow \begin{bmatrix} \text{segs} \\ \text{UR-LIST } \langle \boxed{1}, \boxed{3} \rangle \\ \text{SR-LIST } \langle \boxed{2}, \boxed{4} \rangle \\ \text{FIRST} \begin{bmatrix} \text{simple} \\ \text{UR } \boxed{1} \\ \text{SR } \boxed{2} \end{bmatrix} \\ \text{REST} \begin{bmatrix} \text{segs} \\ \text{UR-LIST } \boxed{3} \\ \text{SR-LIST } \boxed{4} \end{bmatrix} \end{bmatrix}$$

$$(7b) \begin{bmatrix} segs \\ FIRST \text{ simple} \\ REST \text{ e-list} \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ UR-LIST \langle \boxed{1} \rangle \\ SR-LIST \langle \boxed{2} \rangle \\ FIRST \begin{bmatrix} simple \\ UR \boxed{1} \\ SR \boxed{2} \end{bmatrix} \end{bmatrix}$$

$$(8a) \begin{bmatrix} segs \\ FIRST \text{ complex} \\ REST \text{ segs} \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ UR-LIST \langle \boxed{1}, \boxed{3} \rangle \\ SR-LIST \langle \boxed{2}, \boxed{4} \rangle \\ FIRST \begin{bmatrix} complex \\ UR \boxed{1} \\ SR | SR-LIST \boxed{2} \end{bmatrix} \\ REST \begin{bmatrix} segs \\ UR-LIST \boxed{3} \\ SR-LIST \boxed{4} \end{bmatrix} \end{bmatrix}$$

$$(8b) \begin{bmatrix} segs \\ FIRST \text{ complex} \\ REST \text{ e-list} \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ UR-LIST \langle \boxed{1} \rangle \\ SR-LIST \langle \boxed{2} \rangle \\ FIRST \begin{bmatrix} complex \\ UR \boxed{1} \\ SR | SR-LIST \boxed{2} \end{bmatrix} \end{bmatrix}$$

$$(9a) \begin{bmatrix} segs \\ FIRST \text{ empty} \\ REST \text{ segs} \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ UR-LIST \langle \boxed{1}, \boxed{2} \rangle \\ SR-LIST \langle \boxed{3} \rangle \\ FIRST \begin{bmatrix} empty \\ UR \boxed{1} \end{bmatrix} \\ REST \begin{bmatrix} segs \\ UR-LIST \boxed{2} \\ SR-LIST \boxed{3} \end{bmatrix} \end{bmatrix}$$

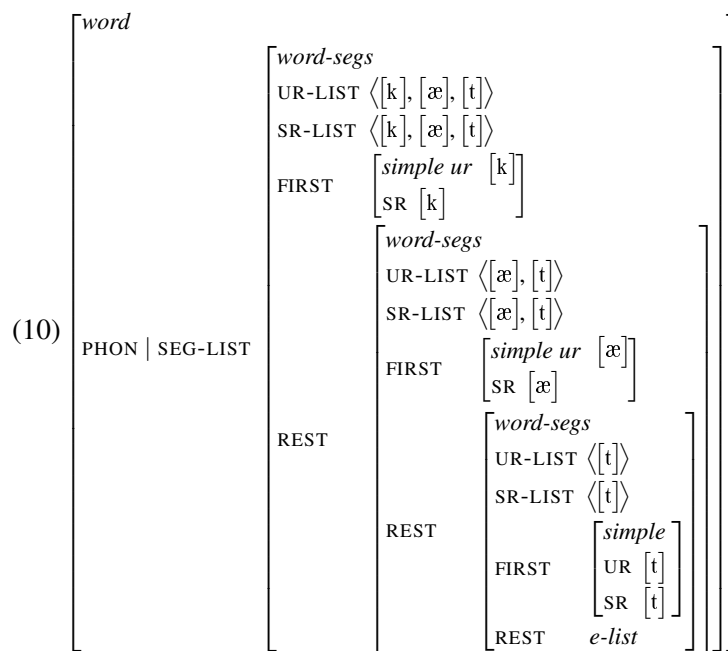
$$(9b) \begin{bmatrix} segs \\ FIRST \text{ empty} \\ REST \text{ e-list} \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ UR-LIST \langle \boxed{1} \rangle \\ SR-LIST \langle \rangle \\ FIRST \begin{bmatrix} empty \\ UR \boxed{1} \end{bmatrix} \end{bmatrix}$$

While it would be possible to group UR-LIST and SR-LIST into a single list, thus simplifying the system, dividing them has an advantage visible particularly in computational implementations: the UR-LIST and SR-LIST objects contain the clearest linear phonetic representation of a particular utterance, phrase or word, which can be invoked to generate the entire structure for the word's PHON. For example, specifying a word's SR-LIST is enough to predict structures for all the possible lexical

items with that surface representation, in a manner in which specifying PHON is traditionally used.

The UR-LIST and SR-LIST features are among the more important ones in this framework: they are used to coordinate generated structures, most importantly appending daughter phonologies to the PHON of mother objects: words into phrases, etc. Using them, rather than simple concatenation of entire PHON structures, allows for using different subtypes of *segs* for different syntactic objects and restricting rules to various levels of the sentence structure. While the PHON structures of words and phrases can be composed of different objects, the core phonological representations are required to be the same. Such an approach combines principles (a) and (c) mentioned in the beginning of the second section.

With these general foundations of the framework in mind, below is an exemplary PHON structure provided according to my proposals for the English word "cat":



As seen above, the *segs* hierarchy is introduced as a feature of SEG-LIST ("segment list"), and not PHON directly. While SEG-LIST is used for linear phonology (and rules operating on segments), other structures can be introduced into the framework, eg. SYL-LIST used for syllables, similar to the solutions introduced in Bird and Klein (1994). This expansion, though possible, will not be covered in this paper.

As can be also seen, the subtype of a *segs* object used above is *word-segs*. The reason for the division of *segs* into various subtypes is to allow formulating rules pertaining to a particular level of representations, as mentioned before. To provide a short example, in English - in much simplified terms - we can postulate a constraint, working on the level of the word, demanding that all surface obstruent clusters have to agree in voicing (in actuality, that would be true for English only for word-final clusters). We restrict this constraint to the word level by evoking *word-segs*:

$$(11) \left[\begin{array}{l} \text{word-segs} \\ \text{FIRST} \mid \text{SR} \quad \text{obs} \\ \text{REST} \mid \text{FIRST} \mid \text{SR} \quad \text{obs} \end{array} \right] \rightarrow \left[\begin{array}{l} \text{word-segs} \\ \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \\ \text{REST} \mid \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \end{array} \right]$$

With such a constraint formulated, the form [kæts] will be well-formed, while [kætz] will violate the constraint. However, because the context is restricted to the word level, the phrase [kæts, doʊnt, flai] is fine, even though the cluster [sd] occurs across the word boundary in the *phrase* object's PHON:SEG-LIST:SR-LIST.

2.3 Final Remarks

The system proposed here is an approach to HPSG phonology in which segments are described dually in terms of their underlying and surface features, and phonological phenomena are handled through constraints restricting or relating the two. Specifying the underlying representation in lexical items would allow us further to leave the surface representation entirely unspecified, thus the possible transformations of the phonemes (such as the ones in example (1)) can be handled purely through constraints.

The simplest constraint linking the surface representation to the underlying one would demand that the SR be identical with the UR. This would, of course, only work in all contexts for an ideal language with no phonological rules (dream on?):

$$(12) \text{ph-str} \rightarrow \left[\begin{array}{l} \text{ph-str} \\ \text{UR} \quad \boxed{1} \\ \text{SR} \quad \boxed{1} \end{array} \right]$$

Actual applications of the UR/SR distinctions will be demonstrated in the following sections using more specific examples, primarily the aforementioned Polish voicing phenomena.

Here I would like to remark that, to focus on the general ideas of this approach, I will not go into the topic of representing individual phonemes in terms of fea-

tures, and my examples will be only as complex as necessary. For the exhaustive analysis for representing phonemes, consult Bird and Klein (1994), Bird (1995, ch. 4) and Höhle (1999). For example, in my analysis, VOICE will be the feature of *rep* (representation) directly, without introducing divisions such as LARYNGEAL/SUPRALARYNGEAL.

3 Word Final Obstruent Devoicing Meets Obstruent Voice Assimilation

The analysis in this section is based around the data and processes in (2), with the goal of adequately describing Polish obstruent voicing processes through HPSG constraints. As mentioned before, there are three elements of the process:

1. Obstruents before other obstruents, including across word boundaries, assimilate their voice to that of the following obstruent, regressively (obstruent clusters have to agree in voicing).
2. Obstruents before sonorants, but not across word boundaries, retain their underlying, distinctive voice.
3. Word-finally, voiced obstruents become voiceless before sonorants or a pause (all word-final obstruents must be voiceless before sonorants or a pause).

The above is true for mainstream Polish, but in south-western variants, the voicing context may be different (Höhle 1999). This will not be dealt with here, although the provided example may easily be altered to account for different voicing phenomena.

The phenomenon of word-final devoicing (3) seemingly acts at the level of the word. However, because it can be "overridden" by voice assimilation (1), the two processes are intertwined and both have to be dealt with at the level of the utterance (thus, the objects *utterance* and *utterance-segs* will be involved).

Before dealing with the rule on the global level, we can use the examples in (2a) as evidence for this phenomenon. We can establish the underlying structure of /d/, in simplified terms, as the following (the voiced form it takes in the intervocalic position):

$$(13) \left[\begin{array}{c} rep \\ UR \left[\begin{array}{cc} obs \\ \text{MANNER} & obstruent \\ \text{PLACE} & coronal \\ \text{VOICE} & + \end{array} \right] \end{array} \right]$$

While the feature MANNER is specified here as *obstruent*, later on I will use the object *obs* to stand for a group of *rep* objects with MANNER: *obstruent*. Höhle (1999), in fact, uses the manner of articulation as a basis for subtypes of segments, eliminating the feature MANNER, while indeed rendering *obs* an existing object.

As mentioned before, because the surface representation can be generated through constraints, declaring anything about the surface structure of the phoneme would be superfluous. Using the same descriptive structure for the underlying and the surface representation, in terms of features (as opposed to alternating between phonemic sorts, cf. Höhle 1999, 3.25) allows us to work with, for any phonological phenomenon, only the features in question, and also organise phonemes into natural groups by invoking the features defining the group. While /d/ in Polish can undergo alternations of its place and manner of articulation (potentially becoming /dʒ/ or, in its devoiced version, /tɕ/), I am going to focus solely on the phenomenon of voicing here, and so, only the feature VOICE will be relevant.

Also note that in this example, it is assumed that all representations are *simple* objects. While epenthesis and deletion exist in Polish, they are, again, beyond the scope of this example and the related constraints can be easily produced by comparing both the global constraints described in the previous section and the examples from the following one.

In Polish, we are dealing with a situation where sometimes the VOICE values of obstruents are determined by the phonological context (before other obstruents, at the end of the phonological word), and sometimes by faithfulness to the UR (everywhere else).

First, we can attempt to translate the rule of voice assimilation:

$$(14) \left[\begin{array}{c} \text{utterance-segs} \\ \text{FIRST} \mid \text{SR} \quad \text{obs} \\ \text{REST} \mid \text{FIRST} \mid \text{SR} \quad \text{obs} \end{array} \right] \vee \left[\begin{array}{c} \text{utterance-segs} \\ \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \\ \text{REST} \mid \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \end{array} \right]$$

This works like the English example previously used - any segment consisting of two obstruents found in an utterance must have a singular value of *voice* in the SR.

Now we have to formulate a constraint determining the voice of an obstruent in any other case (before a sonorant or a pause). We can, for example, formulate the following:

$$(15) \left[\begin{array}{c} \text{utterance-segs} \\ \text{FIRST} \left[\begin{array}{c} \text{simple} \\ \text{SR} \quad \text{obs} \end{array} \right] \\ \text{REST} \quad \neg [\text{UR} \quad \text{obs}] \end{array} \right] \rightarrow \left[\begin{array}{c} \text{utterance-segs} \\ \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \\ \text{FIRST} \mid \text{SR} \mid \text{VOICE} \quad \boxed{1} \end{array} \right]$$

The above constraint demands that if any obstruent followed by a non-obstruent is found in an utterance, its surface voicing has to be equivalent to its underlying voicing. This is, however, imprecise for Polish, since the actual situation is like that only in the word-final context. But neither can we restrict the context to word-final in any *word-segs*, because that would demand every single final obstruent to be voiceless, even if the following obstruent is voiced, which would lead to violating the previous formulated constraint (14).

One way of solving the situation would be to expand the structure of a phoneme with a class of morphology-related features, eg. determining if it is a word-final segment or not. This new information is introduced into the *simple* object as a feature NP (for "non-phonetic"):

$$(16) \begin{bmatrix} \text{simple} \\ \text{UR} & \text{rep} \\ \text{SR} & \text{rep} \\ \text{NP} \mid \text{WORD-FINAL} & \text{boolean} \end{bmatrix}$$

With this expansion, we can determine all word-final obstruents in an utterance. We first need to introduce rules to determine the value of feature WORD-FINAL:

$$(17a) \begin{bmatrix} \text{word-segs} \\ \text{FIRST} & \text{simple} \\ \text{REST} & \text{e-list} \end{bmatrix} \rightarrow \begin{bmatrix} \text{word-segs} \\ \text{FIRST} \mid \text{NP} \mid \text{WORD-FINAL} & + \end{bmatrix}$$

$$(17b) \begin{bmatrix} \text{word-segs} \\ \text{FIRST} & \text{simple} \\ \text{REST} & \text{segs} \end{bmatrix} \rightarrow \begin{bmatrix} \text{word-segs} \\ \text{FIRST} \mid \text{NP} \mid \text{WORD-FINAL} & - \end{bmatrix}$$

Now, we can restate the previous obstruent devoicing rule to include the new information about word-final segments, but operate at the utterance level:

$$(18a) \begin{bmatrix} \text{utterance-segs} \\ \text{FIRST} & \begin{bmatrix} \text{simple} \\ \text{SR} & \text{obs} \\ \text{NP} \mid \text{WORD-FINAL} & + \end{bmatrix} \\ \text{REST} & \neg [\text{UR} \text{ obs}] \end{bmatrix} \rightarrow \begin{bmatrix} \text{utterance-segs} \\ \text{FIRST} \mid \text{SR} \mid \text{LG} \mid \text{VOICE} & - \end{bmatrix}$$

$$(18b) \begin{bmatrix} \text{utterance-segs} \\ \text{FIRST} & \begin{bmatrix} \text{simple} \\ \text{SR} & \text{obs} \\ \text{NP} \mid \text{WORD-FINAL} & - \end{bmatrix} \\ \text{REST} & \neg [\text{UR} \text{ obs}] \end{bmatrix} \rightarrow \begin{bmatrix} \text{utterance-segs} \\ \text{FIRST} \mid \text{UR} \mid \text{LG} \mid \text{VOICE} & \boxed{1} \\ \text{FIRST} \mid \text{SR} \mid \text{LG} \mid \text{VOICE} & \boxed{1} \end{bmatrix}$$

In this case, any word-final obstruent before a non-obstruent in a complete utterance is devoiced, and any non-final obstruent retains its underlying value of VOICE.

This way also, the value of the final obstruent in the cluster is always predicted, allowing the preceding obstruents to assimilate their voice to it in order to retain the constraint (14).

Do note that we can establish other morphology-related features in NP, like a feature STEM-FINAL. In the above examples, the features are encoded into *simple* structure, and would have to operate differently when introduced into a *complex* object, though the above case should suffice for the presented example at least.

4 The Issue of Opacity in Turkish

The purpose of this section is twofold: to return to the processes of epenthesis and deletion, and to demonstrate how accounting for underlying forms can be used to deal with opacity-related issues (rule interaction at more than one level), usually problematic in monostratal frameworks.

In Turkish (based on data for the OT analysis in Sanders 2003, 5.3), two separate processes occur: 1. consonant clusters are broken through epenthesis, and 2. /k/ is deleted intervocalically when a suffix beginning with a vowel follows. However, in the case where /k/ is followed by a consonant, the two take effect at the same time: /k/ is deleted, but triggers epenthesis nonetheless:

1. /baʃ/ + /m/ → [baʃum]
2. /ajak/ + /u/ → [ajau]

But: 3. /ajak/ + /m/ → [ajaum]

The above description is somehow simplified - there are lexical exceptions to this rule, and the /k/ may not completely disappear (in some contexts it may become another consonant, most importantly /j/ before front vowels, or its deletion may lengthen the preceding vowel) - this also depends on dialectal variation. However, here, the case of complete disappearance is assumed, mainly to account for destructive processes in which ghost segments are undesired.

One way to account for such a process in HPSG is to invoke morphology, which is indeed a viable solution. However, I will attempt to demonstrate that with the notion of the underlying representation, HPSG can handle such cases of opacity purely through phonological constraints.

In a typical monostratal framework, introducing the constraints prohibiting both consonant clusters and [k] before vowels could, possibly, lead to a situation where

neither [ajakum] nor [ajakm] are considered well-formed while the form [ajauum] is, but, first of all, we would have no way to arrive at it, and, more importantly, where any cluster of two vowels would have to be acceptable, while in Turkish, that is not exactly the case - the vowel clusters, aside from borrowings, emerge almost uniquely from the deletion of /k/. To account for this fact, the framework would have to postulate the presence, but not articulation, of /k/, as a ghost segment in the cases where it is deleted, but still present for the purpose of epenthesis.

Within the framework presented, it is possible to eliminate the need for such non-surfacing phonemes by translating the two rules (epenthesis and k-deletion) to involve two different levels of representation. This should not be confused with "ordering" the rules, as the two constraints apply simultaneously, but take into account the UR and the SR separately.

To begin with, epenthesis can be formulated by stating that a consonant (here, for simplification, just /m/) can surface either as a single phone, or as a segment:

$$(19) \begin{bmatrix} ph-str \\ UR \quad [m] \end{bmatrix} \rightarrow \begin{bmatrix} simple \\ SR \quad [m] \end{bmatrix} \vee \begin{bmatrix} complex \\ SR \quad \begin{bmatrix} complex-str \\ FIRST \quad [m] \\ REST \quad [m] \end{bmatrix} \end{bmatrix}$$

Now, because epenthesis takes into account the underlying structure of the word, we can translate this rule into the framework by formulating a constraint demanding that any underlying consonant must be followed by a surface vowel:

$$(20) \begin{bmatrix} segs \\ FIRST \mid UR \quad cons \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ REST \quad e-list \end{bmatrix} \vee \begin{bmatrix} segs \\ REST \mid FIRST \mid SR \quad vow \end{bmatrix} \vee \begin{bmatrix} segs \\ REST \mid FIRST \mid SR \mid FIRST \quad cons \end{bmatrix}$$

The context for k-Deletion is the occurrence of a following vowel (in simplified terms, again). Therefore, the necessary constraint would demand that any underlying /k/ followed by a vowel must not surface, ie. be an *empty* subtype *segs* object, appending nothing to the SR-LIST:

$$(21) \begin{bmatrix} segs \\ FIRST \mid UR \quad [k] \\ REST \mid FIRST \mid SR \quad vowel \end{bmatrix} \vee \begin{bmatrix} segs \\ FIRST \mid UR \quad [k] \\ REST \mid FIRST \mid SR \mid FIRST \quad vowel \end{bmatrix} \rightarrow \begin{bmatrix} segs \\ FIRST \quad empty \end{bmatrix}$$

The final result of the two constraints (20 & 21) is that the only permissible situation is the one where deletion and epenthesis co-occur. The presence of an underlying /k/ triggers an epenthesis, but the /k/ does not surface itself, because it is

followed by a surface vowel. The presence of a non-surfacing /k/ may also be used to formulate a constraint demanding the aforementioned vowel lengthening.

5 The Conclusion

With the presented examples and the description, I hope to have shown that it is possible to have a functional phonological framework utilising underlying forms in HPSG, which would tackle neutralisation and opacity without going into arbitrary complexity. Although other proposals for handling phonology in HPSG exist and, indeed, are constantly being developed, the approach presented here aims to be widely applicable and resolve phonetic alternations on purely phonological grounds, while still leaving a lot of space for expansion and not being detached from the structures of morphology and syntax.

References

- Bird, Steven. 1995. *Computational Phonology: A Constraint-Based Approach*. Studies in Natural Language Processing, Cambridge.
- Bird, Steven and Klein, Ewan. 1994. Phonological Analysis in Typed Feature Systems. *Computational Linguistics* 20, 455–491.
- Bonami, Olivier and Boyé, Gilles. 2006. Deriving Inflectional Irregularity. In Stefan Müller (ed.), *The Proceedings of the 13th International Conference on Head-Driven Phrase Structure Grammar*, pages 361–380, Stanford: CSLI Publications.
- Höhle, Tilman N. 1999. An Architecture for Phonology. In Robert D. Borsley and Adam Przepiórkowski (eds.), *Slavic in Head-Driven Phrase Structure Grammar*, pages 61–90, Stanford, CA.
- Prince, Alan and Smolensky, Paul. 1993/2004. *Optimality Theory: Constraint Interaction in Generative Grammar*. Oxford: Blackwell.
- Rubach, Jerzy. 1982. *Analysis of Phonological Structures*. Studies in Natural Language Processing, Warsaw: Państwowe Wydawnictwo Naukowe.
- Sanders, Nathaniel R. 2003. *Opacity and Sound Change in the Polish Lexicon*. Ph.D.thesis, University of California, Santa Cruz.
- Shoun, Abby. 2005. Interpreting and generating formal grammars: A study of HPSG phonology. Honors Thesis: University of North Carolina.