

Conceptual Atomism

Manuel Bremer

§1 Conceptual atomism (CA) is the claim that many if not most concepts cannot be decomposed into a set of conceptual parts or features thus that this set of features is not just necessary but also provides a sufficient analysis of the concept, thus that the conjunction of the features is equivalent to the concept in question. Of course, there are lots of concepts that are derived compositionally from these atomic concepts, and these derived concepts can, obviously, be decomposed again. Still the majority of concepts – and especially those concepts that are in other semantic theories considered to be decomposable – are not decomposable to the conceptual atomist.

I will assume that conceptual atomism is more or the less right. I will not rehearse the arguments for conceptual atomism here (cf. Fodor 1998, 1998a), but rather focus on some of its main ideas and virtues. [They will be taken up below.] Fodor's main arguments for CA are mostly negative anyway: other theories do not deliver what they promise: there are neither enough good definitions around for theories of lexical decomposition nor can theories of prototypicality provide compositionality. CA by definition does not bother about the absence of definitions, and it delivers compositionality for complex concepts, the existence of which is, of course, not denied, and for sentences as well.

Chomsky applied the ‘Poverty of the Stimulus’ argument also to semantics (cf. Chomsky 1993: 21-30). In a similar fashion as in syntax the stimulus presented to children is not rich enough to yield in so short a time the semantic system that children possess:

[T]he growth of the lexicon must be inner-directed, to a substantial extent. ... Barring miracles, this means that the concepts must be essentially available prior to experience, in something like their full intricacy. Children must be basically acquiring labels for concepts they already have, ... (Chomsky 1991: 29).

Chomsky explicitly agrees with Fodor’s theory of concepts on several occasions. For him conceptual atomism is a solution to ‘Plato’s problem’ of accounting for the already given knowledge of language that we possess (cf. Chomsky 1986).

§2 CA is tied to RTM and the idea of a *Language of Thought* (LoT, cf. Fodor 1975). Its model of representation thus inherits their strength. It ties them with ideas stemming from externalist semantics and information content theories. As part of RTM it still, however, keeps the idea of some Fregean sense (i.e. mode of presentation), since the LoT symbols themselves are of distinct syntactic type and thus even in case of co-referentiality present their content in a specific mode.

Important for the relation of theories of concepts and theories of lexical knowledge is the distinction between several layers of representation, which is present in each kind of RTM view on cognitive architecture (cf. Pylyshyn 1984):

<u>Computational</u> Level	<u>Conscious (linguistic)</u> intentional non-conscious $\uparrow$ LOT_1 $\downarrow$	<u>Intentional</u> explanations, cognitively penetrable	Cognitive Levels (in a narrow sense)
<u>Functional</u> <u>Architecture</u> (Design Stance)	transducer information processing • partially genetically fixed • compiled programs (LOT_2)	only <u>functional</u> explanations, cognitively impenetrable	
<u>Implementation</u> (Physical Stance)	φ_1, φ_2 implement ψ_1, ψ_2 if and only if $\varphi_1 \text{ causes } \varphi_2 \equiv \psi_1 \vdash \psi_2$	only <u>causal</u> explanations	basis for supervenient descriptions

FIGURE 1
COGNITIVE ARCHITECTURE

In this picture several levels have to be kept distinct. There is consciousness, which is intentional and often expressed in (natural) language, but not all intentional states are conscious states. Below the level of intentional states, which are the object of intentional explanations, there is a level of pure information processing including *transduction* (i.e. the conversion of sensory stimuli information into LoT-readable representations). The language of thought symbols occur both on the level of intentional mental states as well as receiving the output of the transducers. The deeper levels of information processing (e.g. in phonetic

decoding) are not *cognitively penetrable* (i.e. subject to influences by our beliefs and other propositional attitudes). Some of the capabilities of the lower levels of functional architecture are genetically provided (i.e. innate). On this level of processing there are also the cognitive equivalents of *compiled* programs (using some lower level of the LoT as *machine language*). Other information processing procedures may also be characterized as programs, but rather of a type equivalent to a higher programming language like *Java* (i.e. really containing symbols for the operations to be carried out). Compiled programs are not cognitively penetrable, and may have a role to play, for instance, in some principles of universal grammar (cf. Pylyshyn 1991). On both these levels we have functional explanations. With respect to taking the design stance towards functional architecture (e.g. looking at the workings of our perceptual apparatus) we have only functional explanations, no intentionality is involved here. Computational or intentional explanations are functional explanations in virtue of the representations contained in the explained states, and the semantics of these states. Explanations concerning the functional architecture concern the function of some module or part of the architecture viewed from the design stance towards the whole cognitive system. These explanations although they are functional with respect to the task of some module or transducer need not involve the inner workings of these modules or parts; neither do they refer to any involved LoT-representations apart from those delivered on the output side of the transducers. These levels are the cognitive levels in the narrow sense as they typically involve RTM-explanations involving the LoT. Descriptions in term of LoT, of course, are supervenient descriptions. Below that level we have a level where both the LoT structures as well as the functional architecture are implemented. This is the level of neuro-physiological structures. The reduction problem concerns the question

whether type-type or at least token-token reductions between functionally identified cognitive states/events and physiologically identified brain states can be established.¹

Concepts are (mostly) to be placed in the field of LoT₁ mental processing. Atomic concepts are typically not ambiguous – at least not as far as our natural laws go – and build up non-ambiguous thoughts. The LoT in contradistinction to natural languages, the explicit imagination of which constitutes only *part* of our conscious thought, therefore is the vehicle of thought (cf. Fodor 1998a: 63-72).

Procedures of justifications are to be placed into the intentional field and are typically accessible to consciousness. Not all ways to evidence the presence of a property are conscious. Not all ways of recognizing some property need to be propositional procedures (for example detecting similarities to a stereotype).

CA fits nicely in with theories of animal cognition and the evolution of language. If concepts are pre-given they can explain the learning of a language, and they can explain what we share with other animals. It places the language faculty (in the narrow sense of containing mainly the computational principles, rules and

¹ These questions do not concern us here. Functionalists like Fodor, Chomsky and Pylyshyn have often stressed that although the mind supposedly is implemented in the brain – and we may thus speak rather of the system as the ‘mind/brain’ – the prospects for reductions are rather dim in view of the current state of the neuro-sciences and the nature of their respective language and explanation models. Functionalist theories of the language faculty are far more established, comprehensive and corroborated than the language related results of the neuro-sciences. A future neuro-science might provide, according to Chomsky (2005), a basis for a reduction, but only if it undergoes a conceptual revolution itself, like in the case of the reduction of chemistry to 20th but not 19th century physics. On a more fine-grained analysis there are, of course, more levels inside the broad levels outlined here (cf. Lycan 1987: 37-48). In as much as biological concepts are themselves – in distinction to physical concepts – teleological/functional (e.g. concerning the function of the parts of a cell) even some structures on the neuro-biological level can have a functional description *in that sense* (cf. Churchland/Seinowski 1992). The *stance* terminology stems from Dennet (1987).

representational structures that go into internal word processing, thus excluding articulation...) besides the conceptual system (cf. Jackendoff 1997):

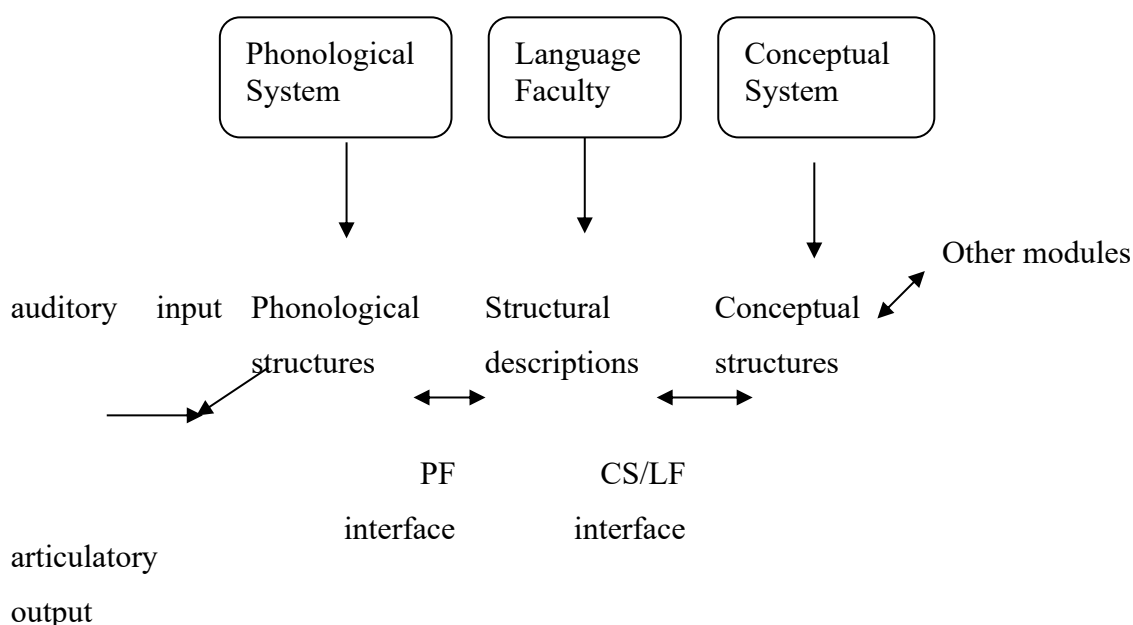


FIGURE 2
THE LANGUAGE FACULTY AND ITS INTERFACES

The structural descriptions, which are the output of the computational system of the language faculty, are interpreted at the interfaces with the phonetical system and the conceptual system (cf. Chomsky 1995, Jackendoff 1997). A derivation ‘crashes’ if some feature in such a structure is not phonetically interpretable. The language faculty operates at two interfaces. It receives input from the phonological system going back to auditory input. It commands indirectly some articulatory output to utter the derived sentence. The derived structures are interpreted at the interface to the conceptual system (sometimes called “logical form” (LF), where logical form in the abstract sense can be shared by linguistic formulae and LoT-representations). The conceptual system and the conceptual

structures it generates are linked to other modules of cognition, which either use these structures or provide content themselves, for example by transducing sensory input. Some animals (like the primates and so-called “higher vertebrates”) may share a large part of our conceptual system. Their interaction with the environment therefore resembles ours. What they lack is the characteristic trait of *homo sapiens*: the language faculty. For reasons of a process of co-evolution of brain structures and language they also lack a human like phonological system (cf. Deacon 1998).

§3 The main criticism of CA mobilizes some version of the ‘incredulous stare’ that most concepts (even REFRIGERATOR) should be atomic and innate. This, however, mellows down to an over-interpretation and some simplistic ways of expressing CA (for the purpose of getting attention ...). The point, of course, cannot be that REFRIGERATOR itself is innate (as some LoT-type) but the means of atomistically tying up to refrigerators. In fact, artefacts might be a bad example, since entities or corresponding theoretical concepts that are explicitly defined are not excluded by CA. CA just claims that more concepts are atomic than one thought of before. DOG might be a better case in point. Discussing the concept of DOORKNOB Fodor (1998: 136-37) himself stresses that having some concept means *resonating to some property* in the appropriate fashion. Even if the physical structure DOORKNOB is composite the concept *might* be unanalysable into jointly sufficient semantic features. Therefore, DOORKNOB might be primitive to us, but this does not mean that the LoT-symbol for DOORKNOB itself is innate. What are innate are the cognitive path ways and mechanisms that introduce us to the concept DOORKNOB in the presence of DOORKNOB:

all that needs to be innate is the sensorium. ... [T]he kind of nativism about DOORKNOB that an informational atomist has to put up with is perhaps not one of *concepts* but of *mechanisms*. (ibid p.142).

That sounds far less eccentric than Assyrians running around with ACCELERATOR. That a concept cannot be defined and thus has to be taken as atomic and thus – at least in principle – as isolated from supposedly closely related concepts – related in which way, anyhow? – does not mean that in the normal process of acquisition a concept comes alone. Given objective relations (i.e. relations pertaining in reality) between the properties referred to, and given the probable sensitivities of the cognitive mechanisms which link a concept to its referent, we may suppose that having these mechanisms in place and interacting with an environment which comes with a structure results in acquiring concepts more or the less at the same time, or, at least, acquiring a concept and a couple of related concepts referring to (metaphysically or nomologically) related properties. Thus DOOR may be not a ‘feature’ present because of a complete decomposition of DOORKNOB, but DOOR and DOORKNOB may be – under normal circumstances – be co-present with DOORKNOB and DOORKNOB, because doorknobs are – under normal circumstances – on doors. A mind/brain able to link to doorknobs is likely to be linked to doors as well. And so for many other concepts.

The co-presence of DOOR and DOORKNOB can be explained this way. It does not cut in favour of some semantic theory. A question arising with this, however, is what else one may expect to be present once both concepts are present in a conceptual system. There should be some analytical links between them.

In contrast to DOOR and DOORKNOB it is not beyond imagination that one can hook up – for some time at least – to SCAR without having WOUND, which is

unfavourable to a semantic analysis involving historic knowledge for some concepts.

§4 Chomsky (1986) distinguishes between accounts of internal language ('I-language') and external language ('E-language'), and he himself considers only the first to be the proper objects of a science of the language faculty. Within the language faculty, on this account, we not only find the core computational system (in the sense of syntactical derivational machinery), but also semantics. This semantics Chomsky (2005) calls 'internal' semantics or even 'syntax', now in a broader sense. It is not semantic in the sense of E-language as it does not concern – at that level of representation – how lexical items are tied up to the world. It consists rather of representations expressing how lexical items are composed and linked. The knowledge thus represented and the derivations and computations flowing from it make up *internal semantics*:

[M]uch of the very fruitful inquiry and debate over what is called “the semantics of natural language” will be understood as really about the properties of a certain level of syntactic representation – call it LF – which has the properties developed in model-theoretic semantics, or the theory of LF-movement, or something else, but which belongs to syntax broadly understood – that is, to the study of mental representations and computations – and however suggestive it may be, still leaves untouched the relations of language to some external reality or to other systems of the mind. (Chomsky 1991: 38)

If many concepts are atomic this can be known in a disquotational theory of truth. Having at some level of presentation a representation of this theory is part of semantic knowledge. This representation is internal semantics. Thus Chomsky and Davidson – otherwise paradigm example of accounts of I-language vs. E-language – could agree that Larson and Segal's theory of semantic knowledge

along the lines of a mental representation of a truth theory for some language (Larson/Segal 1995) was on the right track.

This knowledge accompanied with knowledge of conventions of usage, maybe on the lines of Lewis' theory of *conventions* (1969) or Milikan's theory of *coordinative functions* (2005), provides the language user with guidance, with *semantic rules*. The (justificationist) idea that there are semantic rules is quite compatible with CA. Concepts are not constituted by (semantic) rules, but expressing some concept by a specific *word* within some linguistic community requires rules and possibly shared knowledge of them. So, *identifying* the meaning of a word has to consider these rules, which by this are rules of meaning (semantic rules).

§5 In as much as concepts are atomic in CA one does not have to have some other concept to have a concept in question. So, you can have the concept DOG without having the concept ANIMATE, at least in principle. Once, however, you have both concepts there are ties between these concepts because of the *metaphysical* relations between the *properties* that these concepts refer to. So, once you have all these concepts, and express these concepts in your language, then

(1) Dogs are animate.

should have a privileged, more cognitively entrenched, status in comparison to

(2) There are more dogs in cartoons than there are elephants.

Sentence (1) expresses an analytic dependency between the words “dog” and “animate” because of the conceptual tie between DOG and ANIMATE, because of the metaphysical relation between DOG and ANIMATE. Possession of

concepts in CA does not require the presence of analytic dependencies, but the possession of many concepts brings analytic dependencies around. In contrast to inferential role semantics these dependencies are *not constitutive* of the concepts, but supervene on the conceptual and ultimately metaphysical relations. They *express* some aspect of the metaphysical identity of the property referred to, and its metaphysical relations to other properties.

Even in CA some concepts (like BACHELOR) are explicitly defined (accordingly for some words in some language). In this case we have analytic bi-conditionals. For other concepts, even though they are atomic, there may be meaning postulates (in the form of conditionals) which express irrefutable inferences that are allowed by these concepts. These meaning postulates may be part of the lexical entry of a corresponding word or may be kept in a special semantic belief box. Such meaning postulates may also be called ‘analytic dependencies’. They correspond to *conceptual dependencies* in the conceptual system. The conceptual dependencies depend on metaphysical relations between properties. Meaning postulates capture them in language. They single out some inference as *due to meaning*. Consequently, some sentences are true due to meaning, i.e. *analytic*.

The presence of such conceptual and analytic relations is well-established in linguistics (cf. Chomsky 2005). One may view the attempts at semantic decomposition *and their failure* as really providing not definitions but analytic dependencies. If we look at Jackendoff’s work, for example, we see analyses that stress interrelations between word fields, and work like the following:

- (3) x killed y.
- (4) x lifted y.
- (5) x gave z to y.

(6) x persuaded y that p.

seem to share some conceptual structure in that all of them describe a process which results in some state or event because of what x did. So one may claim (cf. Jackendoff (1989) as an analysis (leaving out the tense):

(3') x cause [y die]

(4') x cause [y rise]

(5') x cause [y receive z]

(6') x cause [y come to believe that p]

This analysis achieves two things: On the one hand more complex concepts and lexical entries are *decomposed* into simpler ones. On the other hand one may supply an (conceptual) *ontological framework* that reads off the basic ontological reference from the constituent structure of the *analysandum*. So we may have:

(3'') [_{event} CAUSE([_{entity} x], [_{event} DIE([_{entity} y])))]

(4'') [_{event} CAUSE([_{entity} x], [_{event} RISE([_{entity} y])))]

and so on.

The first observation about this style of analysis concerns the analyzing concepts like GO or CAUSE and the conceptual resources of the ontological framework like EVENT or STATE. Unless one supports a version of radical holism, in which these concepts are defined using each other, these concepts have to be atomic and thus supposedly innate. The controversy between CA and this style of analysis, therefore, focuses *only* on the question how many concepts are atomic. Neither the presence of atomic concepts nor the wider theoretical context need to be controversial.

The main observation about this style of analysis is that the proposed analyses are just not sufficient. Even if

(7) Miguel lifted the packet.

entails

(8) Miguel caused the packet to rise.

that does not mean that they are (logically) equivalent. There are instances of caused RISINGS which are not instances of LIFTINGS. The sinking fleet or the greenhouse effect may cause the sea level to rise, but we would neither say nor think that they lift the sea. So LIFTING has to have further conceptual components, maybe something along the lines of intentional behaviour. *Maybe*, but this hand waiving at completion often gets us not towards a really sufficient conceptual decomposition. The very reason that we have not just one concept RISE but LIFT, ELEVATE, RAISE, PICK UP, PULL UP, HOIST ... is, according to the conceptual atomist, that there are subtle or not so subtle differences between them that resist being spelled out easily. This is not only a problem about the examples (3) – (6), as one may argue that they only claim a left-to-right entailment. Jackendoff (1990: 53) proposes as a lexical entry for “drink” the meaning expressed by “cause a liquid to go into one’s mouth”. So let’s look at

(9) Ruben drank it.

(10) Ruben caused that liquid to go into his mouth.

Now, that description in (10) may apply as well to events of nipping, guzzling or gulping. Even if these are sub-categories of drinking – as our conceptual analyst may interject – one can easily come up with a story where a person causes a liquid to go into his or her mouth, making (10) true, without being said to drink that liquid (stories may vary from swallowing seamen in sexual intercourse to tripping

over some hose watering the flowers in the garden). So (10) does not entail (9).² Surmising this Jackendoff (1990: 53, 253) classifies his definitions as “oversimplifications” or “tentative”. Still, we lack a collection of non-tentative decompositional lexical entries of interest.

As even the conceptual atomist grants, there are some concepts which can be analyzed into necessary and sufficient conditions, but the majority of concepts and lexical items cannot be so analyzed. So what is said about the example above can be said – or so at least I claim – about the majority of examples presented by Jackendoff or some approach like Pustejovsky’s lexical semantics (cf. e.g. Pustejovsky 1991). Pustejovsky, for example, analyses “door” as: an artificial physical object one can pass through, an aperture. That, obviously, is true as well of the entities referred to by “window”, “entrance”, “gate” or what not.

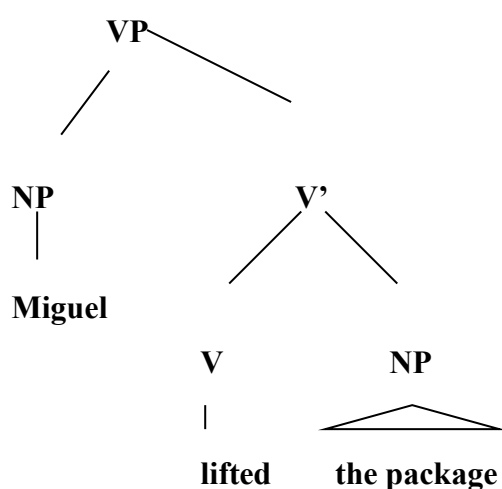
Jackendoff attempts to alleviate these problems by claiming that there are “non-definitional forms of decomposition”, some concepts containing a perceptual part (cf. Jackendoff 2002: 336, 349). Even if there are non-definitional forms of decompositions we still need a theory of how this decomposition works and – given Jackendoff’s logic of constitutive possession conditions – we have the problem: How can concepts such composed be employed and conveyed so that a learner can acquire a new concept/lexical item? That some percept is associated with a concept is probable given the role that paradigms play in concept acquisition and may well be part of a justificatory short route (cf. Bremer 2005: 220-40), all this without being definitional or constitutive. Jackendoff (2002: 387) clearly sees the difficulty of delimiting those items that can be defined from those which cannot.

² Leaving the problem of sentences (9) and (10) having to be used, of course, in situated utterances to anchor the indexical expressions to the side here.

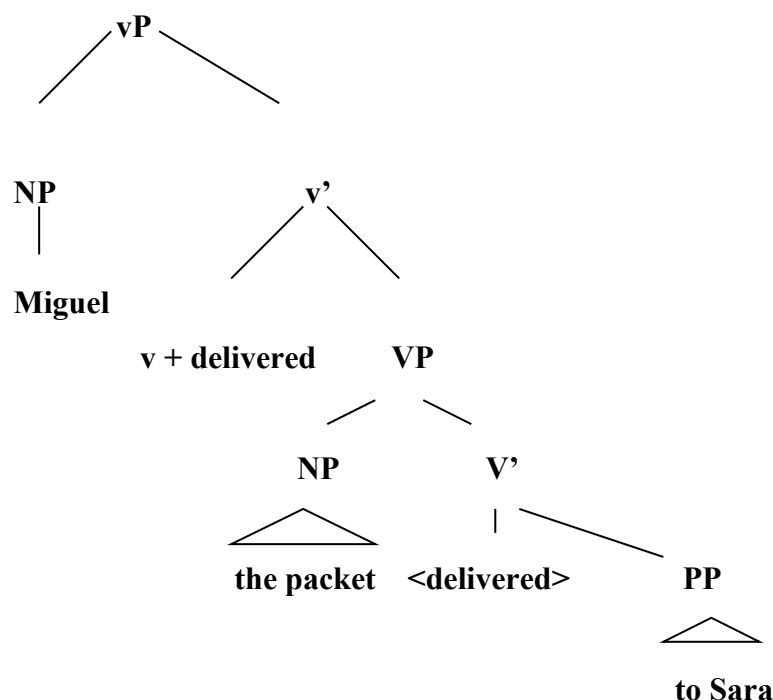
The main observation, however, has one more positive side: the proposals outline and make plausible the existence of analytic dependencies. The form of conceptual analysis practised by Jackendoff or Pustejovsky seems to be able to capture these analytic dependencies.

Further on we see in the examples just considered that some analytic dependencies are central not just to the lexical field in question, but have repercussions in other fields of the theory of language. Jackendoff's point about the ubiquity of analyses with CAUSE is linked to the theory of 'causatives'. The theory of causatives is one of the motivations to introduce into *generative syntax* the so-called 'VP-shell' construction.

Analysing ordinary transitive verbs like "lift" one may come up with a simple tree diagram like



This has some problems of its own, but it certainly cannot be easily extended to di-transitive verbs like "deliver". The standard approach of the VP-shell analysis assumes that there is a constituent type vP above the VP (called "little v", cf. Adger 2003: 121-41). The tree structure – still leaving out tenses and inflection – will then look like



already including *movement* of the verb from V' (leaving a 'trace') to *v*. A constituent structure of this type accounts for some otherwise problematic cases like 'double object construction' and scales nicely up when considering tenses and functional categories. Now, for our discussion of interest is that one of the arguments introducing and motivating 'little *v*' is the causative analysis:

(11) x delivered y to z

(12) x caused [y be at z]

The syntactic relation is assumed to mimic the semantic analysis containing CAUSE, forcing a movement of the verb "into a higher position, which encodes causality" (Adger 2003: 133). The analysis splitting a verb into a causative element 'little *v*' and the remaining 'light verb' leads up then to the 'Hierarchy of Projections', thus is not of minor relevance for the whole enterprise of generative grammar. Decomposition, thus, has a role to play and the presence of

decomposition explains phenomena beyond the narrow semantic field. Good decompositional analyses provide a route to see connections in the language faculty or even cognitive architecture as a whole.

§6 A lexical item has as its *semantic content* some concept with *objective content*. Derivatively then the lexical item has the *objective* content of the concept. The lexical item has more parts than this core semantic content. It also has a part which contains its analytic dependencies. It is linked to its slot in a disquotational truth theory of the language in question. It carries its θ -roles. It also carries features relevant to the core computational system of sentence derivation (cf. Chomsky 1995).

Concepts are the core of the meaning of a word. Other ingredients are pointers to corresponding meaning postulates. The lexical entry may also involve syntactic and pragmatic markers. It involves also pointers to more general principles of meaning (like what it means to take part in a convention and that the overall goal of assertoric utterances is truth). Meaning is conventional in the sense that it is conventional which word is tied to which concept. There are conventions of usage (mostly sticking to uniformly express some concept with some specific word), but use does not constitute meaning. We recognize which word expresses which concept by interpreting usage, but the concept is not constituted by that use. The convention of tying some word to some concept is established as a regularity of usage in some population.

Concepts refer to natural or artificial properties/structures found in reality. Objective relations between these properties (like inclusion, part-whole...) are *metaphysical* relations that are expressed in metaphysical truths. In as much as language wants to capture reality, meaning postulates are incorporated into a

language to mirror such metaphysical truths. Meaning postulates *of this kind* underwrite the *intuition* of analyticity in the sense of theoretical centrality. There are analytic sentences about matters of theoretical centrality, but analyticity does not come down to theoretical centrality.

One criterion concepts like BACHELOR underwrite another group of intuitions of analyticity, namely those of simple analytic sentences like “All bachelors are unmarried”.

Another group of analyticity intuitions stem from thematic role relations. Thematic role postulates are parts of lexical entries since they constrain selection in building sentences.

§7 As there is analyticity there is *a priori* knowledge (for competent speakers of a language). The meaning postulates of a language define what counts as semantic necessity. As there is a analytic-synthetic distinction some problems (occurring in inferential role semantics and some versions of semantic holism) can be kept away. For some concepts and lexical entries – not just one criterion concepts, but also explicitly defined theoretical concepts – there is even a level of internal semantic representation with a non-disquotational but informative theory of truth (i.e. a theory in which the right hand side of a (T)-equivalence does not just use the expression mentioned on the left, but contains its definition).

Informative (T)-equivalences, therefore, lead to *intensions*: We may understand the description of the truth conditions of “p” by another statement even as information about the criteria which justify a usage of “p”. Since we want to know when we have to employ some expression we are interested in the explanatory power of informative (T)-equivalences. Informative (T)-equivalences are the means of *explicit* teaching. The theory now says that “p” varies with some

conditions *q*. If we know that *q* is the case, we, therefore, are *justified* in applying “*p*”. Non-informative (T)-equivalences say little to nothing about the criteria of application of “*p*”. They give us extension. Informative (T)-equivalences can be read as giving us the criteria of justifying employment of “*p*”, i.e. intension.³

[This is, of course, a basic modification of Davidson’s claims and Fodor’s seeming rejection of analyticity.] Davidson says in a later work (1990a: 310):

[A theory of truth] also specifies the conditions under which the utterance of a sentence would be true if it were uttered.

We sometimes possess *criteria* to discriminate the appropriate conditions. These criteria are justification rules for some *other* description of the situation, namely: “*q*”. Informative (T)-equivalences give us meaning (or some approximation to meaning) by logical equivalence. This is still a theory of truth, since the assignment of “is true” to any statement has to employ such criteria.

³ The set of possible truth conditions on the right side of such an equivalence will be far greater than those which can be operationalized as rules of justification. The aim is to specify most simple epistemic circumstances (e.g. perceptions) for complex referents (e.g. a molecular structure). Furthermore, it is not required that a single speaker knows all rules of justification. Putnam’s division of labour in case of reference has to play its role here.

Going back to FIGURE 2 we may give a more detailed picture of the language faculty then:

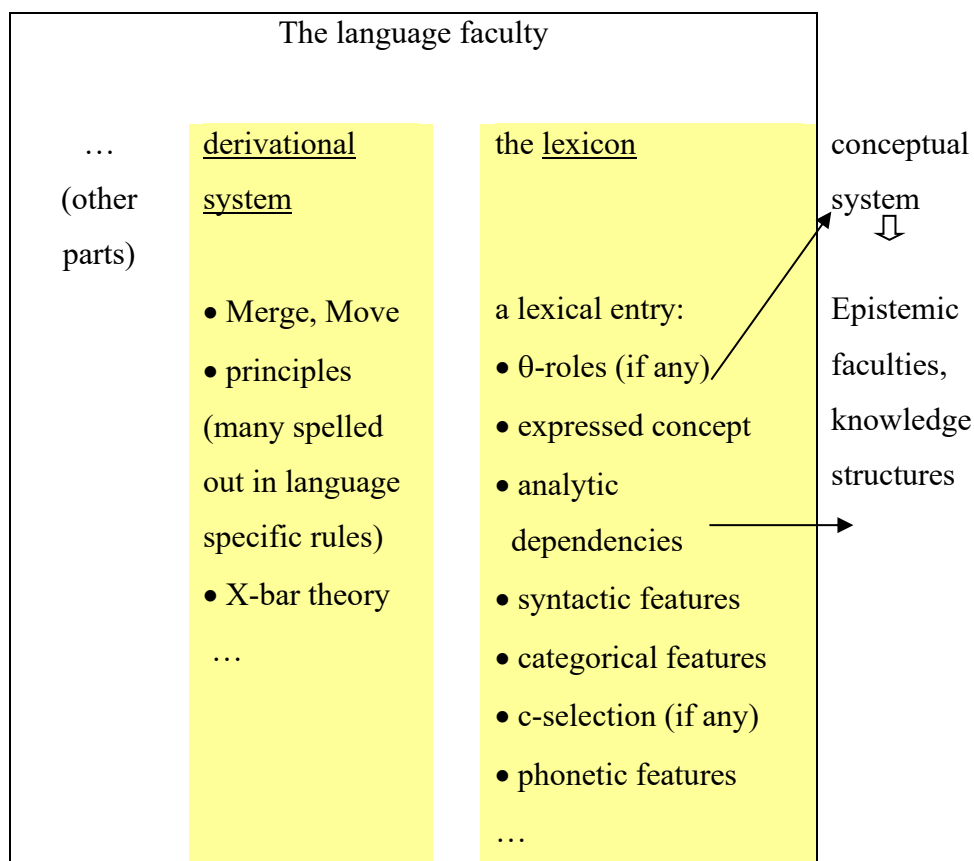


FIGURE 3
SOME INNER STRUCTURE OF THE LANGUAGE FACULTY

In this picture the derivational system, although it uses the lexicon is set apart from the lexicon to highlight the more complex structure of lexical entries. The lexicon not just specifies the θ -roles for words, needed to built structure by Merge, and categorical information on word class, but also those elements related to the discussion here. A lexical entry should contain a link – or a ‘pointer’ as in some programming languages – that connects the lexical item to the expressed concept

in the conceptual system. Part of the lexical entry has to be the set of analytic dependencies. Whether these are stored in the lexicon itself or only linked to by another pointer is of minor importance. In any case these analytic dependencies are related to our epistemic faculties and knowledge structures that exploit these dependencies in the sense of justificationist semantics. The analytic dependencies are linguistic equivalents of conceptual dependencies. Thus the knowledge structures working with them are based on the conceptual system. (Whether there are more parts of the language faculty or more sub-parts of a lexical entry does not concern us here.)

§8 Concepts in the RTM-sense explain the correlation of concepts with properties in the world (as their *objective* content). The objective content of a concept is some property in reality. A concept as a specific *Language of Thought* type (tokened on occasions of thought) also has its syntactic shape and features, which provide it with a mode of presenting the content. A concept in CA is thus more than a mere content connected representation.

Concepts such specified have a *content* and a *mode of presentation* (by being *Language of Thought* symbol types/tokens). CA thus combines the benefits of an informational account of objective content with the Fregean idea of modes of presentations. These modes of presentations occur already at the conceptual level itself. There may be more modes of presentation further up in the level of representation in as much as we know about analytic dependencies. Since analytic dependencies need not be shared by lexical items with the same objective content they acquire a *secondary mode of presentation*, still not a Fregean sense in determining reference, but yet a Fregean sense in being intersubjective.

The account of objective content is externalist and close to theories of information flow (like Dretske's theory (1981)) and accounts of *informational content* in situation semantics (cf. Perry/Israel 1990, 1991, Bremer/Cohnitz 2004). This embeds it in other virtuous accounts of semantic phenomena. Because it also trades in modes of presentations a lot of the criticisms put forth against semantic externalism do not apply!

Routes and procedures of verification are not part of the core meaning of a word then. A lexical entry might be associated with a link to access procedures, but these are not meaning constitutive. One may possess the concept without being able to verify its applications. One can understand the content of a sentence without knowing how to verify or justify it.

Access procedures do not have to have the strength of definitional equivalence. Prototypicality effects have their place here. These procedures are not compositional (as concepts are) and they need not be strict, but only reliable. There are procedures of recognitional ties to concepts, but no recognitional concepts, the conditions of possession of which would be constitutive epistemic conditions.

§9 Most atomic concepts have to be acquired. They are acquired in a reliable way by the corresponding innate mechanisms of connecting the mind/brain to reality. Acquiring the concept DOORKNOB requires acquaintance (transduced sensory causal contact) with DOORKNOB. Hooking the concept up to the property involves being attuned to clues that correspond to the presence of the property. These clues need not be accessible to conscious. The relation between the thus cognitively present clues and the concept nevertheless is *evidential* in the weak sense that the clues significantly raise the probability of the property being

instantiated in the source. One thus learns to use clues. (Again, conditions of possessing concepts are linked to *prima facie* justifying its instantiation by evidence.) This evidential link – to stress this one again – is just motivational: it is not acquiring knowledge about doorknobs or constituting DOORKNOB. The sheer plenitude of different clues working on different occasions excludes them from being constitutive. Recognitional capacities are not constitutive, so that there are no recognitional concepts.

Concept acquisition involves not just causal contact, but involves cognitive mechanisms that lead up to representations (i.e. the concepts) being the consequence of sensory mechanisms. For more complex concepts evidential relations in a stricter sense play a role: If the concept SOCCER is analytically linked to the concept BALL GAME, then being brought to belief by the circumstances that a ball game is going on *raises the probability* that soccer is going on. And a thought containing BALL GAME in conjunction with other thoughts going back to the observed circumstances may *cause* the crucial (observational) belief containing SOCCER. In as much as analytical links tend to be reliable the whole process of acquiring the belief that soccer is played when watching soccer is reliable. Core semantics – in contrast to epistemology as the theory of justification – allows for this dose of foundationalism (cf. Fodor 1987: 118-26).

Here prototypical effects and stereotypes play an important role. Being non-compositional and being barely intersubjectively accessible stereotypes cannot be the concepts, but they can be the typical stepping stone to acquire a concept and recognize the presence of its referent in a situation.⁴ “[H]aving a concept and

⁴ Being nomologically connected to the properties they cannot be *the properties* (i.e. the referents of concepts) since nomological co-instantiation is too weak even for metaphysical identity if natural laws – as is the common conception – are contingent, not to speak of logical identity.

having its stereotype are reliably closely correlated” (Fodor 1998: 138). Our mind/brains work in a fashion that links to properties by intervening structures that show stereotype effects. (In the second example above we may very well identify soccer games by a stereotype of a scene including a pitch and teams, instead of working to SOCCER by way of BALL GAME and what not.)

Identifying properties is in this way mind dependent. We tend to hook up to those properties that we are able to stereotypically identify. We also identify other properties (e.g. in science), but generalizing over experience to stereotypes as indicating some property is the usual experiential way to hook up to some property (i.e. acquiring some concept). This mind dependence, of course, does not make the properties referred to in any sense less ‘real’. The observation concerns our limited access to structures of reality, not some limitation of reality relative to our constructive powers, or whatever. All this applies to artefacts (like doorknobs) and to natural kinds (like water) in the same way:

The kind-constituting property is a hidden essence and you get locked to it via phenomenological properties the having of which is (roughly) nomological necessary and sufficient to instantiate the kind. ... WATER, like DOORKNOB, is typically learned from its instances; but that’s not, of course, because *being water*, is mind-dependent. Rather, it’s because you typically lock to *being water* via its superficial signs; and in point of nomological necessity, water samples are the only things around in which those superficial signs inhere. (Fodor 1998: 156)

Prototypes are thus a major road in applying concepts. The ways of knowing of the presence of doorknobs or identifying doorknobs need not be shared to a degree which allowed speaking of an intersubjective ‘sense’ or ‘mode of givenness’ in the Fregian tradition:

Wittgenstein was right, after all, that the primary check on whether we mean the same by our words is agreement in judgements, but agreement in judgements proves

nothing about agreement in the methods of identifying used in making those judgements. (Millikan 2005:71)

Judgements in situations of usage given the presence of the topic of the description are usually uncontroversial by hidden resources of agreement. The experimental method in science manages to put us into situations where we would have some thought if the prediction was true (cf. Fodor 1994: 92-99). We trust our senses as they are reliable indicators of the properties which the concepts that occur in our observational beliefs are locked up to. That is half of the truth.

Apart from malfunction concepts being triggered by the presence of their referents does not supersede the quest for truth and justification. The major innovation coming with human concepts being part of a compositionally complex representational system is *displaced reference*, or, with linguistic symbols, *displaced speech*. Thinking or talking about what is not obvious or present enables deliberate planning and recollected history. Concepts occur in such displaced thought and talk because of their inferential connections. Present beliefs and one's background knowledge force upon us – with different strength of comparable plausibility, maybe – other beliefs as conclusions and predictions. Tracing concepts such employed from their situation of usage to their situation of justification ('verification'), which may, given enough background theory, *be* the present situation, requires procedures of justification or short routes of *prima facie* credibility instalment on those beliefs.

For members of a language community linguistic expressions can become, because of their meaning and the corresponding conventions of usage, the source to have some belief (thereby tokening some concepts) being in the audience of an utterance.

§10 Even if concepts are either clear as they are atomic or as they are defined concepts are not lexical items. Already the naming convention for concepts suggests the idea of one property: one concept. In proportion to the limited sensibilities of the mechanisms that link concepts to their referents there is *some* indeterminacy of reference (say if the mechanism cannot distinguish between some nomologically co-instantiated properties, and one's property theory counts them still as two properties. The larger problem might be shifted to lexical items. A lexical item – in the ideal case – has one concept as its meaning.

Nothing excludes, however, that one lexical item has more than one concept as its meaning. In cases of clear ambiguity the lexical item can be treated as ambiguous, with the two readings known and separable; there could as well be two lexical entries in this case.

In other cases there may be distinct but in their content (i.e. the properties referred to) overlapping concepts, where we do not recognize that we are dealing with two concepts. In many applications of the lexical item the differences may not count or surface, but in others they may well do so. A – so to say – well-behaved realm of concepts does not guarantee a well-behaved lexicon.

How could the cognitive mechanism that links concepts to linguistic expressions work? The most simple idea might look like a stamp engine:

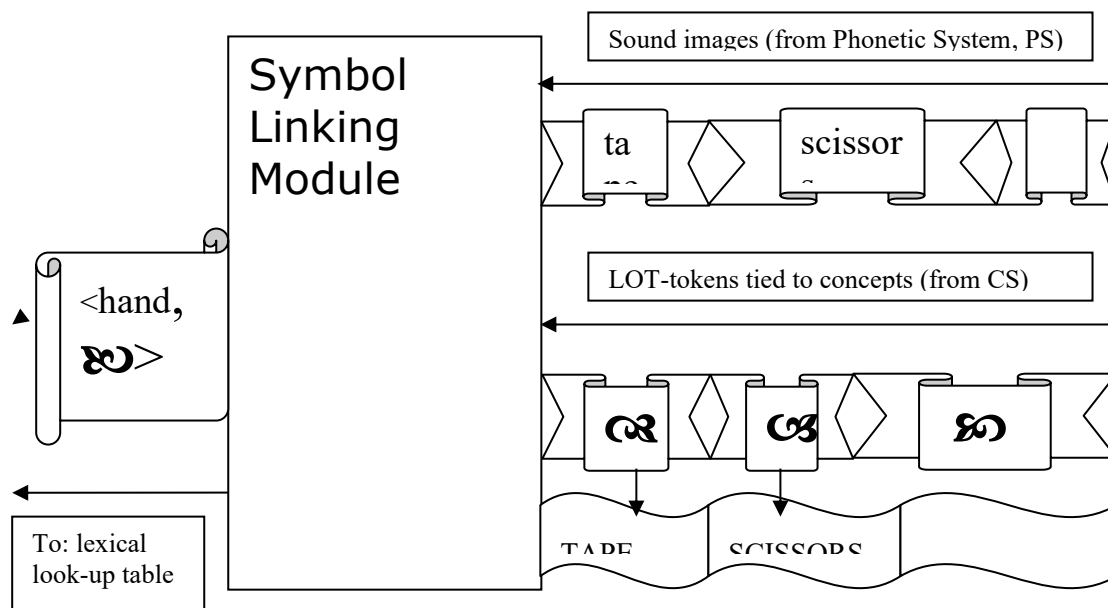


FIGURE 4
NAÏVE PICTURE OF SYMBOL LINKING

There is an input belt to the labelling box; in each cell of the input belt sits a token of a LoT symbol type; this token is labelled with a token of a linguistic expression and the pair is stored into a look-up table. Now, if a token of the LoT type is to be expressed or an expression to be decoded the look-up table is consulted. If that was the case the ideal of an injective mapping between LoT and language (and a bijection between coded LoT types and lexical items) would hold. Ambiguity would only occur because of malfunctioning: either stamping two LoT tokens at the same time or multiply assigning a lexical item to several LoT tokens. How the labelling mechanism works may be considered as *the soft symbol link problem*. Dealing with it one has to provide a theory that combines an account of the mechanism as just outlined from the perspective of the *individual* working

cognitive system with an account of conventions and/or coordinative functions, which play a major role in selecting some word (type) as expressing a concept within a language *community*. This problem includes some tricky questions of social coordination, it seems, but may well be put from the perspective of the intersecting idiolects of individual speakers (i.e. be put from the individual's perspective). In any case the *lexical* assignment problem might be considered to be not that crucial for a theory *of concepts*.

Things – in the most literal sense – may not be that simple. There are properties out there. These properties stand in relations, some of which are mereological. Thus, properties can be combined, whether this combination is natural in the sense of spatial or causal contiguity or connection or not. The thus built complex is metaphysically real, simply said: a property itself. Thus in the absence of any guarantee that our cognition-world alignment mechanism are tuned in to natural properties (in a sense of “natural” to be specified in a substantial metaphysical theory) we may have tuned in to such a property and so have developed a concept that as every concept traces some property out there, but trace a property that as well could be spilt up in two more natural properties (again in a substantial sense to be specified). In this case the conundrums presented by some lexical items could go back to such gerrymandered concepts. Can we ever distinguish a case where two clear concepts are tied to one lexical item from a case where a gerrymandered concept got lexicalized? And is that an important distinction in the first place? If it is not an important distinction in the first place, what does this tell us about the benefits of a well-behaved conceptual system? The plethora of these questions may be considered *the hard symbol link problem*.

References

- Adger, David (2003). *Core Syntax*. A Minimalist Approach. Oxford.
- Barwise, John/Seligman, Jerry (1997). *Information Flow*. The Logic of Distributed System. Cambridge.
- Bendall, Kent (1979). "Negation as a Sign of Negative Judgement". *Notre Dame Journal of Formal Logic*, XX, pp. 68–76.
- Bremer, Manuel (2005). *Philosophische Semantik*. Frankfurt a.M.
- Bremer, Manuel/Cohnitz, Daniel (2004). *Information and Information Flow*. Frankfurt.
- Chomsky, Noam (1986). *Knowledge of Language*. New York.
- (1991). "Linguistics and Cognitive Science: Problems and Mysteries", in: Kasher, A. (Ed.) *The Chomskyan Turn*. Cambridge/MA, pp.26-55.
- (1993). *Language and Thought*. Rhode Island/London.
- (1995). *The Minimalist Program*. Cambridge/MA.
- (2005). *New Horizons in the Study of Language and Mind*. New Edition. Cambridge/MA.
- Davidson, Donald (1982). "Empirical Content", *Grazer Philosophische Studien*, pp.471-489.
- (1982a). "Rational Animals", *Dialectica*, 36, pp. 318-27.
- (1983). "A Coherence Theory of Truth", in: Henrich, D. (Ed.) *Kant oder Hegel*, Stuttgart.
- (1984). *Inquiries into Truth and Interpretation*. Oxford.
- (1990). "Meaning, Truth and Evidence", in: *Perspectives on Quine*, ed. by R. Barret/R. Gibson, Oxford, pp.68-79.
- (1990a). "The Structure and Content of Truth", *The Journal of Philosophy*, pp.279-328.
- (1999). "The Emergence of Thought", *Erkenntnis*, 51, pp.7-17.
- Deacon, Terrence (1998). *The Symbolic Species*. The Co-evolution of Language and the Brain. New York.
- Dennett, Daniel (1987). *The Intentional Stance*. Cambridge/MA.
- (1991). *Consciousness Explained*. London.
- Devlin, Keith (1991). *Logic and Information*. Cambridge/MA.
- Dretske, Fred (1981). *Knowledge and the Flow of Information*. Cambridge/MA.
- (1988). *Explaining Behavior*. Reasons in a World of Causes. Cambridge/MA.
- Fodor, Jerry (1975). *The Language of Thought*. New York.
- (1987). *Psychosemantics*. Cambridge/MA.
- (1994). *The Elm and the Expert*. Cambridge/MA.
- (1998). *Concepts*. Where Cognitive Science Went Wrong. Oxford.

REFERENCES

- (1998a). *In Critical Condition. Polemical Essays on Cognitive Science and the Philosophy of Mind.* Cambridge/MA.
- (2000). *The Mind Doesn't Work that Way. The Scope and Limits of Computational Psychology.* Cambridge/MA.
- Jackendoff, Ray (1983). *Semantics and Cognition.* Cambridge/MA.
- (1989). "What is a concept, that a person may grasp it?", *Mind and Language*, 4, pp. 68-102.
- (1990). *Semantic Structures.* Cambridge/MA.
- (1997). *The Architecture of the Language Faculty.* Cambridge/MA.
- (2002). *Foundations of Language. Brain, Meaning, Grammar, Evolution.* Oxford.
- Kripke, Saul (1982). *Wittgenstein on Rule-Following and Private Language.* Oxford.
- Larson, Richard/Segal, Gabriel (1995). *Knowledge of Meaning. An Introduction to Semantic Theory.* Cambridge/MA et al.
- Lewis, David (1969). *Convention. A Philosophical Study.* Cambridge.
- Margolis, Eric (1998). "How to acquire a concept", *Mind and Language*, 13, pp. 347-69.
- Millikan, Ruth (2004). *Varieties of Meaning.* Cambridge/MA.
- (2005). *Language: A Biological Model.* Oxford.
- Montague, Richard (1976). *Formal Philosophy.* New Haven, 2nd Edition.
- Perry, John (2007). "Situating Semantics: A Response", in: O'Rourke, M./Washington, C. (Eds.) *Situating Semantics. Essays on the Philosophy of John Perry.* Cambridge/MA, pp.507-76.
- /Israel, David (1990). "What is information?", in: Hanson, P (Ed.) *Information, Language and Cognition,* Vancouver, pp.1-19.
- (1991). "Information and Architecture", in: Barwise, J. et al. (Eds.) *Situation Theory and Its Applications*, Vol. 2. Stanford.
- Pustejovsky, James (1991). "The Generative Lexicon", *Computational Linguistics*, 17, pp.409-41.
- Pylyshyn, Zenon (1984). *Computation and Cognition. Towards a Foundation for Cognitive Science.* Cambridge/MA.
- (1991). "Rules and Representations: Chomsky and Representational Realism", in: Kasher, A. (Ed.) *The Chomskyan Turn.* Cambridge/MA, pp.231-51.
- Searle, John (1969). *Speech Acts.* Cambridge.
- Searle, John/Vanderveken, Daniel (1985). *Foundations of Illocutionary Logic.* Cambridge et al.
- Stein, Edward (1996). *Without Good Reason.* Oxford.
- Vanderveken, Daniel (1991). *Meaning and Speech Acts. Formal Semantics of Success and Satisfaction.* Cambridge.