

*Logic in Pragmatics**

Hiroyuki Uchida

Abstract

This paper argues against the complete elimination of logical introduction rules from the pragmatic inference system. To maintain the consistency of the inference system as a whole, which is meant to support one's truth-based judgment over propositions, the inference system should have access to both introduction and elimination rules. I show that the inclusion of introduction rules in the pragmatic inference system neither overgenerates propositions expressed nor cause non-terminating inference steps.

1 Free enrichment and alleged overgeneration

According to Relevance Theory (RT for short, Sperber & Wilson 1986/95), reference assignment and disambiguation are not the only pragmatic processes involved in the recovery of propositions expressed by (or the truth-conditional content of) utterances.

- (1) a. Every presenter [*in the pragmatics session of CamLING07*] was impressive.
- b. John took out a key and opened the door [*with the key*]. Cf. Hall (2006)

RT assumes that, given the linguistically provided information outside the square brackets in (1), the hearer can pragmatically add the contents given in the brackets when she recovers the proposition expressed (in appropriate contexts).¹

* I am grateful to Robyn Carston, Nicholas Allott and Alison Hall for reading (part of) drafts and making suggestions. They do not necessarily agree with the final version and all the mistakes and misconceptions in this paper are the author's.

¹ More accurately, given the linguistic meaning outside the square brackets in (1) which the hearer can recover by decoding the semantic information encoded with the language expressions that the speaker used in the utterance, and given the context in which the utterance was made (which includes the speaker's intention), the hearer can pragmatically enrich to the truth-conditional content of the utterance which includes the pragmatically added contents inside the square brackets in (1). Crucially, the hearer may enrich the meaning of an expression before she recovers the encoded meanings of the other expressions. For more details about free enrichment, see Carston (2002).

Stanley (2002) claims that this free enrichment overgenerates. Suppose the sentence in (2) is uttered in a context in which the propositions in (3) are available as contextual premises.² Then, according to Stanley, RT wrongly predicts that the meaning of (2) can be enriched by conjoining it with the contextual proposition (3), deriving the proposition expressed in (4). Stanley himself does not literally identify this process as &Introduction, but for the purpose of this paper, let me identify this process as &I applied to (2) and (3a).³ The recovery of (4) in this way would give the hearer enough cognitive effect, because an application of MPP between another contextual premise (3b) and (4) would lead to the relevant conclusion *John will not live long* (= *R*).

- (2) John smokes. (= *P*)
- (3) a. John drinks. (= *Q*)
b. If John smokes and drinks, he will not live long. (= $(P \& Q) \rightarrow R$)
- (4) John smokes [*and drinks*]. (= $P \& Q$)

Addressing this criticism, Hall (2006: 95-96) follows Sperber & Wilson (1986/1995) and postulates Conjunctive Modus Ponens (Ponendo) Ponens (CMPP) as in (5). With CMPP, one can derive the relevant conclusion, *John will not live long* (= *R*), without applying &I.

- (5) Conjunctive Modus Ponens:

1. $(P \& Q) \rightarrow R$	Premise 1
2. P	Premise 2
3. $Q \rightarrow R$	1, 2, CMPP
4. Q	Premise 3
5. R	1, 2, 4, MPP

In her solution to the alleged overgeneration problem, Hall suggests a weaker claim that because of CMPP, the hearer does not have to use &I in order to derive the relevant conclusion *John will not live long*. Thus, the hearer *can* derive this

² My example sentences are overloaded with different kinds of information. They may represent sentences uttered, linguistic meanings or propositions, where propositions may be trivial, expressed or contextual (though in my views, trivial 'propositions' do not acquire propositional status until they are recognized as propositions expressed (but then they stop being trivial anymore). For example, (2) may represent the sentence uttered, but when it is identified with *P*, it represents a proposition.

³ I make this assumption because the main argument of this paper is that inclusion of &I in the spontaneous inference system at the basic level neither leads to overgeneration with enrichment nor leads to infinite inferences. Reviewing what Stanley really meant in his overgeneration claim and arguing against it is not part of this paper's aim.

conclusion as in (5), without deriving the undesirable (4) as the proposition expressed by (2).

However, with this weaker claim, to prohibit the derivation of (4) as the truth condition of (2) in *any* instance of interpreting (2) in context, one would need some additional explanation why the hearer *always* uses the inference steps as in (5) rather than the application of &I followed by MPP, when free enrichment is involved.⁴

In this paper, I argue that introduction rules can be used in pragmatic inferences in general. Thus, after showing that it is problematic to eliminate the &Introduction rule from the pragmatic inference system in section 2, I provide an explanation about why &I is not used in enrichment in section 3, though the pragmatic inference system itself is equipped with this rule. I also argue that CMPP is only a convenient shorthand for a particular combination of inference steps, rather than an actual inference rule defined over logical connectives.

The stronger interpretation of RT's proposal⁵ is that spontaneous inferences do not use (logical) introduction rules at all (and thus, (5) is the only way of deriving the conclusion *R*, given (2)~(3)). The reason for postulating this stronger hypothesis is not only the alleged overgeneration of propositions expressed by way of free enrichment. Sperber & Wilson, among others, argue that spontaneous inference should not have access to introduction rules because, otherwise, the system would generate infinite or non-terminating inferences. In section 4, I briefly explain this infinity problem and then show that the problem is not caused by the use of introduction rules in the system, and thus eliminating &I or other introduction rules is not the right way of coping with this problem. Section 5 shows some proofs to support my arguments. Section 6 deals with some loose ends and comments about use of logic in pragmatics from a general viewpoint. Section 7 provides concluding remarks.

This paper is based on certain theoretical assumptions. When we say that an inference system is incomplete with regard to the intended semantics, the 'intended semantics' does not mean the semantics of the inference data that the system aims to explain. It means the system-internal semantics that the person who proposes the system must define or provide for the language representations that are manipulated

⁴ For some explanation that Hall suggests about why CMPP is preferred to &I followed by MPP, see Hall (2006: 96).

⁵ Sperber & Wilson explicitly rule out introduction rules from the inference system. Thus, strictly speaking, what I call the stronger claim is the only possible interpretation of their proposal. On the other hand, they carefully avoid any modification of the logical system which may potentially underlie their inference system. Thus, it is theoretically possible to interpret their claim as the weaker one, in which we can use & with restrictions for purely application reasons, such as efficiency of inferential processes, even if that is not what Sperber & Wilson had in mind as a possibility.

in the system. When a deductive system is proposed, both the syntactic rules that define the well-formed syntactic objects (i.e. propositions in propositional logic) and the syntactic inference rules which operate over those syntactic objects (such as MPP, &I, &E etc.) must be presented, but that is not enough. It must also be specified how those syntactic objects and inference rules are intended to be interpreted. Moreover, such intended ‘denotation’ of syntactic objects and rules must be modelled in a well-defined semantic structure, such as the Boolean lattice for the classical propositional logic.⁶ The ‘completeness’ of an inference system with regard to the intended semantic is a matter of checking (or ‘proving’) whether the syntax and the semantics match up completely in the proposed system, where both of these are system-internal concepts. Let me elaborate a little on this point with informal schemas.

- (6) a. Syntax: $\{\dots, \phi_1, \dots, \phi_n, \dots\} \vdash_{\text{CLP}} \{\dots, \varphi, \dots\}$
 b. Syntax simplified: $\phi_1, \dots, \phi_n \vdash_{\text{CLP}} \varphi$
 c. Semantics: $\|\phi_1\|^M, \dots, \|\phi_n\|^M \models \|\varphi\|^M$

In classical propositional logic, the syntax derives a set of propositions from a set of propositions, as shown in (6a) (CLP abbreviates Classical Propositional Logic). For the sake of simplicity, however, let me discuss the syntactic derivability as if we derived a proposition (rather than a set of propositions) from a set of propositions, by getting rid of ‘irrelevant propositions’ (indicated by ... in (6a)) which do not play an essential role in the inference that we discuss at each stage.⁷ That is, as shown in (6b), the proposition φ is syntactically derivable from propositions $\phi_1, \dots, \phi_n$ (I also omit the set notation, $\{\cdot\}$, in the Antecedent to the left of $\vdash$ as well). This syntactic derivation is solely dependent on the set of syntactic

⁶ The word ‘denotation’ may be misleading to linguists, because linguists tend to assume that denotations are some concrete objects in the world, but that is not necessarily the case. Functions from possible worlds to sets of individuals, for example, might be denotations of some logical expressions (such as predicates). Whatever one assigns as the (intended) interpretations of syntactic objects count as ‘denotations.’

⁷ As we see in section 4, a problem of the ‘impure’ definition of $\vee I$ is that as well as it introduces the truth functional connective, it recovers one of these ‘irrelevant propositions’ from the background (i.e., $\{\dots\}$) and put it in a noticeable place. This latter operation is structural weakening in the Succedent and because structural weakening exists independently $\vee I$, solving the (infinity) problem by controlling the application of (impure) $\vee I$ is not only in the wrong track, it does not completely solve the problem, either, as we see in section 4. The same applies to &I in the Antecedent side with regard to the ‘impure’ left rule in Gentzen sequent presentation, as in the inference from $p \vdash p$ to $p \& q \vdash p$, which has incorporated structural weakening in the Antecedent. In contrast, the pure &I in (16a) abstracts away from the structural weakening. See section 5.

derivation rules which include &I, &Elimination, MPP ($= \rightarrow$ Elimination) etc. Now, in order for the system to be used in application, defining such syntactic rules is not enough. We have to define the semantic interpretation rule that interpret both the syntactic objects (such as $\phi_1, \dots, \phi_n, \phi$) and the syntactic derivability relation $\vdash$. Let $\models$ represent the interpretation of the syntactic derivability $\vdash$. This semantic relation $\models$ is harder to explain without introducing formal details, and I only make an informal presentation.⁸ It informally means that for all the models M in which $\|\phi_1\|^M, \dots, \|\phi_n\|^M$ are all 1 (or True), $\|\phi\|^M$ is also 1. This semantic computation is based on the standard interpretation of the logical connectives, &, $\vee$ and $\rightarrow$ in terms of the truth tables. Or one may use the equivalent truth-condition presentations for the connectives, such as, for all models, M , $\|\phi_1 \& \phi_2\|^M = 1$ if and only if $\|\phi_1\|^M = 1$ and $\|\phi_2\|^M = 1$ as the semantics of &, etc. Now, because this semantic relation $\models$ applies generally, independent of the verdict of the syntax, if we assume that p , q and $p \& q$ are all well-formed formulas in the language that the inference system uses, then the semantics validates the argument, $\|p\|, \|q\| \models \|p \& q\|$ (the reader can check the validity by drawing truth tables, for example), even if we eliminates &I from the syntax and we can no longer ‘syntactically’ derive the sequent $p, q \vdash p \& q$. Thus, the truth-based semantics above would validate an argument that the syntactic system without &I (but which still makes use of the form $\phi_1 \& \phi_2$ as a well-formed formula and which still uses &E and all the other elimination rules) can no longer support. This syntax would then be incomplete with regard to the suggested truth-based semantics.

The main part of this paper is just an elaboration of this incompleteness argument against the syntax without &I, relative to the ‘truth-based’ semantics as was sketched above, but let me concentrate on the formal status of the semantics I have just sketched. As I said above, this semantics is independent of the syntactic rules (that is why some arguments may be validated without the syntax being unable to support them). On the other hand, it is still part of the language system in two ways. First, any logical language without the provision of such a formal semantic structure as the one given above is not complete as a system. Without the intended semantics, the syntactic objects and rules may potentially be interpreted in different ways and thus the proposed syntactic system cannot be rigorously evaluated in terms of what it can do in application. Second, it is the intended semantics that is comparable to the data. The impression that one can compare the syntactic rules

⁸ Wansing (1993), among others, interprets a proposition, p , as the set of information states in which p is true. Then, given additional rules to interpret connectives, such as &, $\vee$ and $\rightarrow$, the semantic relation $\models$ corresponds to the subset relation in set theories. This interpretation creates the Boolean lattice structure as the intended semantic structure.

directly to the inference data is illusory: one gets that impression because one has already assigned some arbitrary (or the ‘most natural’) interpretations to the syntactic rules. Because of these, any theory that makes use of some language in presentation should provide a precise definition of the intended semantics. Whether the intended semantics actually corresponds to the semantics as in the data is a separate issue. If it turns out that some system is incomplete with regard to the intended semantics, then it simply means that the system does not work in a complete way system-internally.

Having said that, because I am mostly concerned about ‘truth-based inferences’ of the spontaneous inference system, and because classical logic (which I claim to be essentially the same as one’s spontaneous inference system at the basic level) is sound and complete with regard to the Boolean semantics as I sketched above (or, more accurately, the Boolean lattice structure as in footnote 8), I implicitly assume that the ‘intended semantics’ of the spontaneous inference system is simply the standard Boolean semantics, where the derivability in the syntactic derivation schema, $\phi_1, \dots, \phi_n \vdash \phi$, is interpreted as the semantic validity argument such that, for all the models in which $\phi_1, \dots, \phi_n$ are true, ϕ is also true, as shown in (6).⁹ This is convenient, because, as I said above, classical logic is generally complete with regard to this ‘truth table’ semantics on the one hand, and the truth-table semantics is in close correspondence to one’s truth-based judgments over propositions in on-line inferences on the other hand. Thus, by using the Boolean semantics as the intended semantics, we can mostly ignore the difference between the intended semantics of the deductive system and the semantics that models the actual interpretation. Because of this, in this paper, I am often careless about the distinction between the intended semantics of the inference system and the semantics of the inference data. In this way, I aim to show that the stronger claim made by RT (as well as some other systems that eliminate introduction rules from the inference system) is problematic from both theoretical and applicational viewpoints.

Finally, I assume that the pragmatic inference system has properties of ‘deductive systems’ at the basic level. In other words, I assume that all the rules in the system, including the syntactic rules, the interpretation rules, and the relation(s) over semantic objects in the intended semantics, apply with their full generality. One might define additional rules in the syntax to explicitly control the application of some syntactic rules (such as &I), but then one would also have to provide the

⁹ I mostly ignore the semantics of sub-sentential expressions in this paper, because the main topic of this paper is truth functional connectives. See footnote (31), though.

intended semantics for those additional rules, to show that the restrictions in terms of those additional rules really work in the intended way in the semantics.¹⁰

2 Problems of eliminating &I from the pragmatic inference system

In this section, I discuss some of the problems of eliminating &I from the spontaneous inference system. First, if truth-based judgment is at least part of one's spontaneous inference ability, then the inference system without &I become incomplete with regard to the intended semantics. Also the system uses a rule (that is, CMPP) which the syntactic rules of the connectives involved in that rule do not support. In other words, the system adds an additional theorem which is not supported by the rules of the logical connectives that are manipulated in the theorem. Thus, the inference system fails to be fully deductive.

- (7) a. $p, q, (p \& q) \rightarrow r \vdash r$
 b. $p \& q, (p \& q) \rightarrow r \vdash r$

Consider the two sequents (or the two 'arguments,' if we see them from a semantic viewpoint) in (7a) and (7b). The two atomic propositions p and q on the one hand and one complex proposition $(p \& q)$ on the other have the same truth-based interpretation in the antecedents of the sequents.¹¹ If the inference system cannot make use of &I, one cannot syntactically explain the same role that these formulas play in truth-based interpretations. Without &I, the syntactic system can still derive the entailment relation from $(p \& q)$ to p , and from $(p \& q)$ to q via &Elimination but that is not complete with regard to the truth-based semantics.¹² Thus the validity of (6a) cannot be explained without &I, and the stronger claim by RT requires CMPP as an essential rule, as has been discussed already.

¹⁰ Note that there is asymmetry between the syntax and the semantics throughout the rule formation in the language system. That is, when we control some rule applications, we must first specify the control in the syntax, and then define the interpretation of such control in the semantics. The control in the syntax must be sound and complete with regard to the intended semantics so that the control can really work in the intended way.

¹¹ I do not show a minimal pair in the other direction. That is, in addition to (7), evaluation of the validities of a pair of sequents such as, i) $p, q, p \rightarrow (q \rightarrow r) \vdash r$ and ii) $p \& q, p \rightarrow (q \rightarrow r) \vdash r$, would be necessary to show the equivalence of the roles of 1) p, q and 2) $p \& q$, in an antecedent of a sequent. I omit such a pair because one can prove them only with MPP and &E.

¹² See section 6.1 for further remarks about my recognition of the truth-based semantics as (part of) the intended interpretation of the inference system. I add some comments about Braine and O'Brien's 'procedural' semantics in section 6.3.

However, the reason why CMPP's successive application of $(p \& q) \rightarrow r$ to p and q separately in (5) does not cause a problem for the inference system as a whole is the logical equivalence relation in (8). The proof of this equivalence requires $\&I$, as well as $\&E$ (elimination).¹³

$$(8) \quad (p \& q) \rightarrow r \dashv\vdash p \rightarrow (q \rightarrow r)$$

Imagine a logical system with $\&E$ but without $\&I$, and call it CPLE. (7a) is not provable in CPLE. Now, imagine that we add CMPP to CPLE and call the resultant system RT. Then, (7a) is provable in RT. Because RT is equipped with MPP (which is an 'elimination' rule of $\rightarrow$ and thus, would be preserved in RT), (7b) is also provable in RT. If we want to maintain congruence (cf. footnote 21) and transitivity of the system, which are both essential for a fully deductive system, then, there should be some path from p , q separately to $p \& q$ as one formula, whereas RT is lacking in this path, that is, $\&I$. Thus, the system fails to be fully deductive.¹⁴ To make my arguments clearer, let me review the inference in (7a), and consider how the stronger claim by RT would relate this inference to the inference in (7b) and how the inference system without $\&I$ would syntactically recognize the semantic equivalence between 1) p , q as separate propositions on the one hand, and 2) $(p \& q)$ as one complex proposition on the other, in the antecedent of a sequent in the truth based semantics.¹⁵ Suppose that the inference system were lacking $\&I$ (as in the stronger interpretation of RT's proposal). Then, the inference system would not recognize (7a) as a derivable/provable sequent via $\&I$. However, suppose that this hypothetical inference system were equipped with CMPP instead. Then, a person using this inference system could tell that (7a) is a derivable sequent. Next, the person equipped with this inference system could compare the inference in (7a) to another inference in (7b) which her inference system can recognize as derivable, too, but this time by using MPP. Now, she notices that the pair p , q and the complex formula $(p \& q)$ are replaceable with each other in the

¹³ See the proofs in (23) in section 5.

¹⁴ In a sense, RT is comparable to a hypothetical Combinatory Categorical Grammar system which insists that they can use function composition without abstraction rule (N.B. function composition as a higher order theorem is derivable from abstraction and association as axioms). I do not investigate whether we can preserve the deductive nature of the inference system without $\&I$ but with CMPP by decomposing CMPP into some axioms in a 'modular' system as is sketched in section 6.1 and 6.3. My guess is that there is not a lot of promise. Controlling structural associativity in a multi-modal deductive system is easy because structural rule neither introduces nor eliminates connectives such as $\&$, $\vee$ and $\rightarrow$. In comparison, restricting the use of $\&I$ with presence of $\&E$ even in one mode would cause a problem to the deductive system. I leave further investigation about this point for another paper.

¹⁵ Again, I show the recognition of the equivalence relation in one direction only.

antecedent of an otherwise equivalent sequent (that is, 7a and 7b are identical except for p , q and $p \& q$) and this replacement does not change the validity of the argument. To the degree that the person using this system finds that this is almost always the case, she can reflectively recognize the semantic fact that p , q on the one hand, and $(p \& q)$ on the other, have the same (truth-based) interpretation in the premise of a sequent.¹⁶ However, the hypothetical inference system that she is equipped with is still lacking a direct way of supporting this semantically valid inference, because it is lacking &I. Instead, the inference system would recognize it indirectly, as I have shown above.

I do not find a convincing reason to explain our intuition about the valid argument in (7a) and its relation to another valid argument in (7b) in this indirect way (or in this reflective way). In an informal (truth-based) semantic inference, the conditional $(p \& q) \rightarrow r$ requires both p and q to be true (as the standard truth table shows) in order for r to be true, but that is exactly how the propositions p and q are interpreted in the premise of an inference, and thus one can semantically conclude that r is true. The rule &I in classical logic is postulated just to support this semantic judgment, and the inference system as a whole should be equipped with it, in order to make the system complete with regard to this semantic inference.

As I explained above, one can see this incompleteness issue in terms of replacement possibilities between (sets of) propositions. Whenever p , q on the one hand, and $(p \& q)$ on the other, appear inside the otherwise identical set of premises, RT can explain why the result of the inferences are the same only in an indirect way. Thus, for all the other cases in which our semantics tells us that the choice between 1) ϕ , φ and 2) $\phi \& \varphi$ in the premises of an argument does not influence its truth-based validity,¹⁷ RT would require some rules analogous to CMPP. For example, consider the semantically valid argument, $p, q, \neg(p \wedge q) \vee r \vdash r$. RT without &I would require another rule analogous to CMPP to support its validity. The fact that $\neg(p \wedge q) \vee r$ and $(p \wedge q) \rightarrow r$ are inter-derivable in classical logic does not help, because, without any introduction rules, the RT inference system cannot recognize them as equivalent (again, without adding another formation rule such as $p \rightarrow q \dashv\vdash \neg p \vee q$ whose addition to the system would further spoil the deductive nature and completeness of the system without introduction rules for truth functional connectives).

¹⁶ This is not always the case, because of the sequents, i) $p, q \not\vdash p \& q$ and ii) $p \& q \vdash p \& q$ in RT's system. But this case is trivial in the current argument, because i) is exactly the sequent that RT claims that one needs to exclude. Again, I argue that i) should be maintained and i) does not do any harm in its application in spontaneous inferences.

¹⁷ Carson (p.c.) claims that it is not clear why this is something that the spontaneous inference system should be expected to explain. See section 6.4 for this point.

In a similar way, RT might need an additional axiom to deal with the following case.¹⁸ Consider the set of premises $P1 = \{p, q, (p \& q) \rightarrow r, \neg r\}$. How would RT without $\&I$ but with CMPP deal with this premise set? As one possibility, RT can first apply CMPP between p and $(p \& q) \rightarrow r$, concluding $q \rightarrow r$, to which RT can apply MPP with q , deriving r . Then, RT would get an inconsistent set of premises $\{r, \neg r\}$ from which RT would conclude a contradiction that is, $\perp$. Thus, following this first route, the inference system would derive, $p, q, (p \& q) \rightarrow r, \neg r \vdash_{RT} \perp$.

Alternatively, starting from the premise set $P1$, RT can first apply MTT between $(p \& q) \rightarrow r$ and $\neg r$, concluding $\neg(p \& q)$. Then the resultant premise set would become $P2 = \{p, q, \neg(p \& q)\}$. Now, in terms of the truth based semantics, the three propositions $p, q, \neg(p \& q)$ cannot be all true in any model. Therefore, $P2$ should semantically lead to a contradiction. However, RT cannot syntactically derive a contradiction from $P2$, because RT is not equipped with $\&I$. Thus, following the second route from $P1$, RT's verdict is $p, q, (p \& q) \rightarrow r, \neg r \not\vdash_{RT} \perp$. Comparing the two inference-routes starting from the same premise-set $P1$, we might argue that the verdicts of the RT's spontaneous inference system is inconsistent, as well as pointing out again the incompleteness of the inference system with regard to the intended semantics (i.e. the verdict of the second route means that there should at least be one semantic model in which all the four formulas $p, q, (p \& q) \rightarrow r, \neg r$ are true, whereas, as the reader can easily check by drawing a truth table, there does not exist such a model).

Allott suggests that maybe RT would get rid of MTT from the inference system, but as far as the treatment of the premise set $P1$ is concerned, I am not sure if that is the route they would take. RT would have to deal with the premise set $P2$ anyway. Thus, RT might define another axiom to derive the sequent, $p, q, \neg(p \& q) \vdash_{RT} \perp$. Again, addition of such an axiom that is not supported by the basic rules for the connectives involved in the axiom would spoil the deductive nature of the inference system.¹⁹

¹⁸ Nicholas Allott (p.c.) suggested this case to me. Though I use the set of premises and the basic line of arguments that Allott provided, my analysis may differ from his.

¹⁹ As I implied in the introduction, I do not have a strong view about the claim that the spontaneous inference system is not deductive at all. If that was the case, most of my arguments against a spontaneous inference system without introduction rules would become irrelevant. However, given the productivity of spontaneous inferences, and also given the more than superficial similarity between the logical systems as are investigated by proof theorists and the inference systems that are investigated by psychologists or more empirically minded linguists/philosophers, I do not think that we have to accept the split between the two types of logical systems at the foundational level, rather than at the level of application.

I have shown some reasons not to eliminate &I from the inference system. Now, should we still preserve CMPP in our spontaneous inference system? If the system is equipped with &I, we do not need CMPP as an inference rule. However, because of the equivalence of the two formulas in (8), one may still use CMPP as a shorthand for a set of inference steps in application. That is, given (8), applying $(p \& q) \rightarrow r$ successively to p and q does not cause any problem to the inference system as a whole, based on the basic property of the logic in which replacement of a sub-formula in a sequent with a logically equivalent formula does not influence the provability of the sequent.²⁰ Analyzing interpretation data is beyond the scope of this paper, but to the degree that data suggest that one may use the inference step as in CMPP, we can still treat CMPP as a shortcut in application.

In this section, I have shown that the inference system cannot recognize the truth-conditional equivalence between certain propositions without &I. Though truth-based arguments are not the only kind of arguments that the spontaneous inference system is intended to support, as long as such semantic arguments are at least important in spontaneous inferences, failing to support them at the basic level of the system compromises the explanatory power of the system as a theoretical tool. In terms of congruence,²¹ RT's stronger claim can cause situations in which

²⁰ Sperber & Wilson (1986/1995: 99-100) argue that the rule of CMPP is psychologically plausible in terms of the maximization of the usefulness of the new information that one gets in spontaneous inferences (see also Hall 2006: 96). Roughly speaking, when one has the proposition in the form of $(P \& Q) \rightarrow R$ as one premise, the possibility of finding P and Q separately as other premises is greater than finding $P \& Q$ together. I reserve my view to this point in terms of probability. In terms of efficiency, however, processing one premise after another makes some sense. To support that point, the proof of the sequent $P, Q, P \rightarrow (Q \rightarrow R) \vdash R$ is algorithmically less complex than the proof of the sequent $P, Q, (P \& Q) \rightarrow R \vdash R$, in terms of the complexity measure based on the number of connectives involved in the proofs (i.e. the latter proof includes the introduction of & to conjoin P and Q which increases the complexity of the proof by one). Thus, if one can automatically interpret $(P \& Q) \rightarrow R$ as $P \rightarrow (Q \rightarrow R)$ during a spontaneous inference given the availability of P at that stage of the inference, then the deductive steps that use CMPP will be less complex than the steps using &I, followed by MPP (N.B. the former inference steps would not really derive $P \rightarrow (Q \rightarrow R)$ from $(P \& Q) \rightarrow R$, rather, given the availability of P , one can apply $(P \& Q) \rightarrow R$ directly to P , whereas the logical equivalence between $(P \& Q) \rightarrow R$ and $P \rightarrow (Q \rightarrow R)$ as shown in (23) justifies this successive application of $(P \& Q) \rightarrow R$ to P and Q . So the inference steps will be exactly like (5), though in our system, with a formal underpinning from the basic logical system). The point is that my analysis is totally compatible with the use of CMPP as an application shortcut (CMPP is a higher order theorem that is provable from the basic axioms), and it can also show the efficiency of the spontaneous inference steps using CMPP.

²¹ Informally, for all X, Y . X is congruent with Y (in the antecedent of a sequent) iff, for all Z_1, Z_2, W . $[(Z_1, X, Z_2 \vdash W) \Leftrightarrow (Z_1, Y, Z_2 \vdash W)]$ (where X, Y, Z_1, Z_2 are meta-variables for sets of propositions, and commas between such sets represent the set-union $\cup$). That is, X and Y are

one may replace a pair of propositions with a (complex) proposition without influencing the provability of the sequent, but in which one cannot derive the latter from the former directly. This corresponds to the incompleteness of the syntax with regard to the Boolean semantics, but seen from a different viewpoint, we could also say that the syntactic system uses a rule (i.e. CMPP) which is not supported by the basic rules of the connectives involved (that is, the rules for $\&$ and $\rightarrow$ but without the introduction rules). This spoils the deductive nature of the inference system.

3 Alleged over-generation by way of enrichment

In this section, I show that RT does not have to eliminate $\&I$ to prevent the alleged overgeneration via free enrichment.

In propositional logic, $\&I$ requires as premises two formulas that can be assigned truth values, as is informally shown in (9).

- (9) a. Syntax: $p, q \vdash_{\&I} p\&q$
 b. Semantics: If $\|p\| = \text{True}$ and $\|q\| = \text{True}$, then it follows that $\|p\&q\| = \text{True}$.

Because of this, if one also assumes that the proposition expressed is the first meaning representation that the hearer derives out of a language expression to which the hearer can assign a truth value (in context),²² it follows that $\&I$ (or introduction rules for any truth functional connectives) cannot be used in the derivation of a proposition expressed.²³ Consider (2)~(4) again. Propositions in (3a, b) as contextual assumptions are fully propositional on their own. On the other hand, the semantic content of the sentence in (2), which is uttered by the speaker, acquires a fully propositional status only after it is recognized as the proposition expressed by that utterance. Thus, one can conjoin (2) with (3a) via $\&I$ only after recognizing the semantic content of (2) on its own²⁴ as the proposition expressed.²⁵

congruent in the antecedent iff for no sequent, replacing one with the other in the antecedent influences the provability of the sequent.

²² I stipulate that the hearer does not assign a truth value to a minimal proposition. This does not prohibit the hearer from evaluating a minimal proposition or a trivial proposition in the recovery of the proposition expressed. See the end of this section.

²³ Sperber & Wilson explicitly make their mental logic operate over some non-propositional representations, so we have to do some more work to apply this criteria to their system. See section 6.1.

²⁴ This is inaccurate because one may enrich the semantic content of (2) before applying $\&I$ or any other rules for truth functional connectives. Either the literal meaning of (2) or the result of enrichment based on it can enter into $\&I$ with (3a). But such enrichment cannot include rules of truth functional connectives because of the semantic requirements of those connectives, plus the assumption that the semantic content of (2) does not become truth evaluable just because the

Assuming that the hearer may derive only one proposition expressed per utterance, it follows that (4) cannot be the proposition expressed by the same utterance of the sentence in (2). Note that in this explanation, one does not need CMPP as an actual logical inference rule. CMPP might still be used to describe an on-line inference step that arises as a result of routinization of certain logical inference steps in application. But use of CMPP does not spoil the fully deductive nature of the inferential system as a whole, either. With &I, the system can recognize the equivalence between the role of the two premises p and q separately, on the one hand, and the role of $(p \& q)$ as one complex premise, on the other.

I stipulated that one can apply introduction rules for truth conditional connectives only after one enriches the meaning of the overtly used expression to the *proposition expressed*. As well as its role in keeping the inference system sound and complete with regard to the intended semantics, the assumption is supported by the semantic claim that nothing that the hearer recovers from the language expression during the process of deriving the proposition expressed needs to enter into truth based inferences (other than the proposition expressed itself). For example, if the hearer does not see the literal meaning of *John drinks* as relevant enough in the context of an utterance, she does not need to assign a truth value to the proposition that corresponds to that literal meaning (and she does not have that proposition enter into truth based inferences). She only has to assign a truth value to the proposition that she takes as being expressed, say, “John drinks alcohol,” for example. Thus, from some sort of economy consideration, I assumed that the speaker does not assign a truth value to a proposition unless she either sees it as the proposition expressed or it is one of the contextually available propositions (which, because of the roles that contextual premises play in inferences, should be assigned a truth value by definition). However, the proposal would have a problem if one had to apply a truth based logical inference rule to a propositional representation that has not yet been accepted as the proposition expressed in other well-attested cases. Some might argue that ‘trivial propositions’ as in (10) are such cases.

semantic content of it is type t expression (or, informally, just because what is uttered is a ‘sentence.’

²⁵ Hall (2006) postulates two kinds of pragmatic inferences, local inferences that are applicable to sub-propositional expressions, and global inferences that apply only to fully propositional expressions. Allott (p.c.) suggests that this division might inherently be present in Sperber & Wilson (1986/1995). One can regard my proposal in section 2 as one interpretation of this division between two kinds of inferences.

- (10) a. John has a brain.
 (uttered to express the proposition, *John is smart*.)
 b. Meg is human.
 (uttered to express the proposition, *Meg may make mistakes, etc.*)

An argument against my proposal above would be that, in order to derive the propositions expressed (e.g., the ones provided in the parentheses in (10a) and (10b)), the hearer has to evaluate the literal meanings of (10a) and (10b) as trivially true propositions. Some might argue that such evaluation involves truth based inferences.

However, (10a) and (10b) do not pose a problem for my proposal. To recognize the literal meanings of (10a) and (10b) as trivially true, one does not have to apply proper logical inference rules. In other words, to recognize them as trivially true, one does not have to have the trivially true propositions interact with other contextual assumptions in terms of logical inference rules.²⁶ To explain this point, I first evaluate an application of &I from semantic viewpoints and then come back to inferences involved in trivial proposition cases. In the case of &I, what the semantics of & conjoins is a pair of truth values. Thus, one must know the truth values of *P* and *Q* in order to compute the truth value of *P*&*Q*. But before *P* and *Q* are accepted as propositions expressed or unless *P* and *Q* are contextual propositions (which are by definition fully propositional), one cannot decide on the truth values of these propositions.

Because semantic validity arguments (e.g., for the provable sequent $P, Q \vdash P \& Q$, the semantic argument will be, *for all the models in which P and Q are true, P&Q is also true*) abstract away from the choice of a model, some might argue that one does not really need to assign truth values to all the premise propositions of an inference rule for a truth functional connective. They might argue that the choice of a model (which will correspond to the context of the utterance in question) is irrelevant even in inferences for truth functional connectives and thus, one does not need to assign truth values to the premise propositions in such inferences. However, quantification over models presuppose the truth evaluability of each proposition in any of those models, and if one cannot see some bit of the decoded meaning as truth evaluable in spontaneous inferences, one may not use that bit as an input to a truth functional inference rule, such as &I. Thus, though there are still some speculative elements in my proposal, I assume that use of proper classical logic inference rules in spontaneous inferences requires the full truth evaluability of

²⁶ Such propositional interactions prior to the recovery of the propositions expressed are not necessary even when the literal meanings of (10a, b) are recognized as being informative enough and are accepted as propositions expressed. Contextual premises, plus the linguistic meanings of the relevant expressions, will provide enough clues without such interactions.

all the premise propositions. In use of &I in a spontaneous inference, one has to see “John has a brain” (for example) as the proposition expressed, before one uses it as a premise of &I.

In comparison, let me examine a case of a trivial proposition in more detail. Nicholas Allott (p.c.) claims that, to see that (a) “John has a brain” is trivially true, one would need to retrieve or construct (b) “John is human” and (c) “Humans have brains” and let the three propositions interact somehow. However, there is a crucial difference between this sort of inference on the one hand, and inferences for truth functional connectives (or any inference rules which have properly corresponding rules in classical logic) on the other. For presentation reasons, let me change the proposition (c) into a different form, that is, into (c′) “For all individual x , if x is human, x has a brain.” In Allott’s example above, (a) is the semantic content of the language expression uttered²⁷ (i.e., *John has a brain*), whereas (b) and (c′) are contextually available propositions which the hearer may use as premises of an inference rule. Because (b) and (c′) are full propositions (i.e. fully truth evaluable) by definition, the hearer may apply MPP between them and conclude (d) “John has a brain.” Because the literal meaning of the utterance, that is, (a), is the same as (d) (in the role that it can play in truth based inferences), the hearer decides not to take (a) as the proposition expressed, and thus does not assign a truth value to (a) (or more accurately, the hearer does not take the proposition “John has a brain” as the proposition expressed by this particular utterance of the sentence *John has a brain*). Note that this process does not require an assignment of a truth value to the semantic content, *John has a brain* which one presumably derives as the result of the linguistic decoding. What is required instead is the recognition that this literal semantic meaning and the contextually derived proposition (d) are the same,²⁸ and

²⁷ One could call “John has a brain” a proposition even though it is judged to be trivial for whatever reasons, rather than the ‘semantic content’ of the sentence, because this proposition exists independently of the English sentence, *John has a brain*. However, because trivial propositions are the cases in which the hearer does not judge such a proposition as the proposition expressed by the utterance (and thus does not accept it as the proposition for the utterance, in my analysis), I do not use the word *proposition* for (a) here, to avoid unnecessary confusion. In order to be complete, I should consider all the different reasons why such propositions are judged to be trivial, relative to the language expressions used in the context. For lack of space, however, I leave such a complete exposition for another paper.

²⁸ To be complete in my arguments, I would need to show case by case that this identification process actually does not involve any use of logical inference rules as in classical logic. Things will be easy if there are only two cases involved in trivial propositions, that is, either (a) is a tautology or (a) as a candidate for the proposition expressed and (d) as a conclusion derivable from contextual assumptions only are exactly the same proposition, because the use of classical logic inference rules is not necessary in the triviality judgment in either of these two cases. But I am not certain whether this is always the case. I will leave it for further research.

thus, identification of (a) as the proposition expressed by this utterance does not allow the hearer to do anything more than she could do otherwise.

Though my informal exposition here still contains some arbitrary elements and should be formulated more accurately in some other occasions, it relates the trivial proposition case as in Allott to cases that involve tautologies, such as *Boys will be boys* or *People who study math study math*. If tautologous propositions are real tautologies, by definition, they are always true and thus the addition of them into the premise set of propositions will not allow the hearer to derive any other conclusions that she could derive without them. Thus, given an utterance of the sentence *People who study math study math*,²⁹ the hearer may start enriching its meaning before she lets the literal meaning of the utterance interact with other positions in the context at all.

Dealing with Allott's example above, I let the semantic content of the sentence uttered interact with another contextually derived proposition. This is not problematic because the analysis only prohibits 'truth based' interaction between the linguistic meanings of the expressions used and contextually available information. Thus, after accepting my assumptions, one can still enrich the meanings of component expressions by using the information provided by contextual assumptions. In fact, one may even 'mimic' some of the seemingly truth based inferences.³⁰ For example, with model theoretic relations such as sub-set relations, one can mimic logical entailment relations without deriving a fully propositional representation. In this way, one can enrich the meanings of predicate expressions *smart* or *human* via set-containment relations, for example, without deriving a full proposition. On the other hand, I argue that proper logical introduction rules do require fully truth evaluable elements as premises (just as is the case in standard logical systems) and thus, they cannot be mimicked in terms of relations between sets.³¹

²⁹ What proposition is expressed by the use of this sentence will depend on each context and is not relevant to the discussion here.

³⁰ Of course, the claim is that such 'seemingly' truth based inferences are not really truth based. In fact, the claim is stronger than that. It would state that pragmatic inferences that modify the meanings of the expressions uttered prior to the recovery of the proposition expressed are not 'propositional' inferences, whether propositions are interpreted in terms of their truth values or not (see section 6.1). But I do not have space to explain this point, and thus I use 'truth based inferences' as the guiding criteria.

³¹ I regard generalized conjunction as in Partee and Rooth (1983) only as a rule of PF-LF mapping. For example, at an intermediate stage of the syntactic derivation for the sentence, *Jack and Eva smoked*, the syntactic system might interpret the natural language expression *and* as a lambda term, $\lambda x.\lambda y.\lambda P.P(x)\&P(y)$, so that this derived functor expression can be applied to the individual terms *jack'* and *eva'* successively, deriving another lambda expression, $\lambda P.P(\text{jack}')\&P(\text{eva}')$. However, note that the logical connective $\&$ itself stays as a truth functional connective in any of these lambda terms (i.e. it conjoins two propositions, such as $P(x)$ and $P(y)$). Thus, the use

Finally, I briefly discuss an example of a minimal proposition that Carston suggested (p.c.). Interpreting the utterance of the sentence *John poisoned Bill and Bill died* in a relevant context, the hearer may recover the minimal proposition “John poisoned Bill and Bill died” (= A) in the process of recovering the proposition expressed, “John poisoned Bill and Bill died because of that poisoning” (=B), adding the causal relation to the truth-conditional content via enrichment. However, nothing in my analysis prohibits this enrichment process. All I am claiming is that this enrichment process is not ‘truth functional’ inferences, as &I, &E and MPP etc. are. Causal relations are not truth functional relation, as one can see from the meaning of *because*, which is not a truth functional connective, as one can see in a sentence such as *John came because Eva came*.³² Evaluation of sub-propositional elements or even propositional elements recovered from the sentence uttered is still possible, and from empirical considerations, is necessary, as Carston suggests. But my argument is that though such inferences are still pragmatic inferences (and thus, are constrained by the principles of Relevance), there is a formal difference between such inferences and ‘truth functional’ inferences which are deriving (sets of) propositions from (sets of) propositions solely based on their truth value assignments (in all models). Given that difference, I stipulate a certain feeding relation between these two kinds of pragmatic operation at the theoretical level. That is, ‘non truth-functional’ processes which includes enrichment theoretically feed into the truth functional ones that include the rules for logical connectives.

I leave for further research exactly which inferences can be mimicked in this way and which cannot be.

This section explained why introduction rules are not applicable in enrichment. The next section deals with the alleged ‘infinite inference’ problem.

4 Alleged Infinity problem caused by introduction rules

This section briefly addresses the claim that if one’s pragmatic inference system were equipped with logical introduction rules, one would run infinite or non-terminating inferences. Because such non-terminating inferences are not attested in

of generalized conjunction in terms of the lambda abstracted terms as above does not influence the fully truth functional status of &, $\vee$, $\rightarrow$, etc at the level of logical forms.

³² The fact that the proposition B entails the proposition A does not mean that B must be deduced from A by using ‘truth functional’ inferences as is used in classical propositional logic. Entailment relations may come from the preservation of the lexical meanings, and at least in this example of minimal proposition, none of the classical logic rules is used in the enrichment process from A to B (because, again, the causal relation is not truth functional, and neither is the temporal precedence relation, such as “After P, Q”).

interpretation data, it must be the case that the pragmatic inference system does not have access to introduction rules. Arguing against such claims, I show that this problem is caused independently of the use of introduction rules and should be solved independently.

Johnson-Laird (1997: 391) claims that introduction rules, if they are used in spontaneous inferences, may lead to infinite inference steps, as schematized in (11).³³

- (11) a. $P, Q \vdash \&I P \& Q \vdash \&I P \& Q \& P \vdash \&I \dots$
 b. $P \vdash \vee I P \vee Q \vdash \vee I P \vee Q \vee R \vdash \vee I \dots$

However, the alleged infinity in (10a) is because of the expansion of ‘ P ’ to ‘ P, P ,’ and ‘ Q ’ to ‘ Q, Q .’³⁴ It is not because of $\&I$ per se. Thus, eliminating $\&I$ from the system does not solve the problem completely, to the degree that the problem exists. Also, with regard to this structural expansion rule, note that one occurrence and more than one occurrence of the same formula have the same interpretation in truth-based inferences. Thus, the alleged infinity might be just a matter of the imperfect representation system, rather than some imperfection of the inference system. In fact, even at the level of represented deductions, logicians have tried to eliminate un-decidability induced by structural rule applications. Without going into details, one may apply a structural rule only when the consequence of that rule application is required by the next step of the inference.³⁵

In (11b), $\vee I$ presupposes weakening of the succedent set. Because the standard introduction rule for $\vee$ implicitly includes the structural weakening in the Succedent side, one has to separate the concept of $\vee I$ and the concept of weakening, first, and then find out which of these has created the alleged infinity problem.³⁶ Because $\vee I$ persists across different logical systems with different

³³ Braine and O’Brien also describe this version of the problem. Cf. O’Brien (2004).

³⁴ This notation is slightly sloppy, because P, Q must stay as premises of inference in order to be interpreted as ‘And.’ Gentzen sequent presentation captures the semantic equivalence of P, Q and $P \& Q$ in the antecedent of a sequent in a better way (see 16a), though comparison is not straightforward. Because of some technical details, (16a) formally corresponds to $\&E$, rather than $\&I$. Such technical details, however, do not matter. With $\wedge L$ and $\wedge R$ in (16a), the Gentzen system is complete with regard to the intended interpretation, whereas the system without $\&I$ (such as RT’s stronger claim) is not.

³⁵ Braine and O’Brien proposed a similar, but a different proposal in spirit. That is, they modify the underlying algorithm of their system. Because they divide rules of inference into groups which are not supported by the underlying logical system, it causes several problems, incompleteness as one. See section 5.3.

³⁶ Došen (1988) and Belnap (1996), among others, recommend rule presentations which separate the two concepts, a) rules of connectives and b) structural property of the system (where

structural management properties (e.g. presence or lack of structural weakening), and because it is weakening in the succedent that increases the number of propositions such as Q and R in (11b),³⁷ the infinity problem, to the degree that it gives problems to the inference system, is a matter of structural weakening rule, rather than $\vee I$. Thus, one cannot fully control this problem merely by eliminating $\vee I$. Just like expansion of the formulas, it is beyond the scope of this paper to discuss whether weakening does cause problems to the spontaneous inference system, and if it does, how to control it. One may adopt Intuitionistic logic which is lacking in weakening in the succedent, for example. But using a sub-structural logic makes the system incomplete with regard to the truth based semantics. Thus, it cannot be used to explain one's truth based inferences in a complete way.³⁸

Instead of modifying the underlying algorithm of the inference system, I would rather control structural rules at the level of application, as was suggested above. That is, one might set up the forward looking inferences in such a way such that one may apply structural weakening rule only if its output is required by a further inference step.³⁹ Informally, this means that one weakens P to P, Q , only if, say, one has $(P \vee Q) \rightarrow R$, as another premise.⁴⁰ Alternatively, one may try a proof representation system which does not incorporate the structural weakening into the rule of $\vee I$, but which leaves the weakening rule implicit, so that the spurious ambiguity that is caused by application or non-application of the weakening rule simply does not arise. This analysis requires some technical explanation, and I leave the details for another paper.⁴¹

As another variety of the alleged infinity associated with introduction rules, some might argue that recursive applications of $\&I$ followed by $\&E$ would produce infinite inference steps, but this infinity does not arise in standard proof representations without a Cut, such as Gentzen sequent presentation without Cut. Some proofs are listed in section 5 (see (17)~(20)).

b) is explicitly represented as structural rules separately from a). In that conception, $\vee I$ is independent of structural weakening in the Succedent. For example, "The rules for the logical operations are never changed: all changes are made in the structural rules." (Došen, 1988: 352).

³⁷ One can rewrite (10b) as $P \vdash P, Q \vdash P, Q, R, \dots$ etc., without introducing $\vee$.

³⁸ Whether Intuitionistic logic is still useful in a 'modular' inference system in one of its modules is a separate issue. See section 6.1.

³⁹ One can prove that controlling structural rule application in this way does not influence provability of sequents. Thus, the system will stay complete. See section 6.3.

⁴⁰ In this case, structural weakening feeds into $\vee I$, which feeds into $\rightarrow$. Thus, in this particular case, it leads to the same result as Braine and O'Brien. But the way that we achieve it is better, for the reason that we explained already. In this proposal, $\vee$ itself is freely applicable, as long as there are P and Q in the Succedent side.

⁴¹ Roughly, notions such as 'monotonicity' and 'purity' may be assigned to the system itself. See Avron (1993) for the explanation of these ideas. See Wansing (1998:92) as well.

Finally, I discuss a more sophisticated infinity argument. Consider (12).

- (12) (Non-) Frame problem.
 a. Antecedent Set $\vdash_{\Delta}$ Succedent Set
 b. $P \vdash_{\{Q\}&I} P \& Q$
 c. cf. $P, Q \vdash_{\&I} P \& Q$

(12a) represents a spontaneous on-line inference step. Though the logical inference rules are the same as in classical logic, (12a) distinguishes between two kinds of databases that are used as premises. The antecedent embodies the set of premise propositions that are active in the context, including the proposition expressed by the utterance. To draw a conclusion in the succedent set, one can also use premise propositions in the ‘dormant database’ set Δ , which contains the whole of the (propositional) knowledge that one has.

With these assumptions, some might argue that the inference system would wrongly predict the existence of an infinite inference as in (13).

$$(13) \quad P \vdash_{\{Q, R, S, \dots\} \&I} P \& Q \vdash_{\{R, S, \dots\} \&I} P \& Q \& R \vdash_{\{S, \dots\} \&I} P \& Q \& R \& S \vdash_{\{\dots\}} \dots$$

In (13), one may extract one proposition after another from the dormant database set and conjoin them with the proposition P in the active premise set. If one assumes that the amount of one’s knowledge is almost infinite, this model wrongly predicts that one may actually run an almost infinite inference.⁴²

However, note that this alleged infinity is not a matter of $\&I$ per se. As I have already pointed out, in the antecedent set, P and Q as separate propositions on the one hand, and $P \& Q$ as a single complex proposition on the other, play the same role in the classical logic. Thus, the above infinity problem will arise independently of the use of $\&I$. What is problematic then is the introduction of Q, R, S into the active data-base, not the conjunction of those newly introduced propositions with a proposition that is already in the active data base. Thus, what one needs is a systematic way of constraining the introduction of propositions from the dormant database to the active database.

This section has shown that use of introduction rules in the inference system is not the cause of the alleged non-terminating inferences, and that the elimination of introduction rules does not solve the problem.

⁴² Also, in a spontaneous inference, one does not typically access all the pieces of knowledge that one has, even if the pieces of knowledge are relevant to the argument one is making.

5 Some Gentzen sequent proofs

In this section, I show that successive applications of &I (or &R in this section) and &E (or &L) do not lead to undecidability. I also show that $(p \& q) \rightarrow r$ and $p \rightarrow (q \rightarrow r)$ are inter-derivable. The proofs here are elementary, and are a simple application of the Gentzen sequent presentation of classical logic as in Girard (1987) or Takeuti (1987).

(14) Sequent to prove (e.g.) $p, q, (p \& q) \rightarrow r \vdash r$

Gentzen sequent proof representation places the sequent to prove at the bottom of the derivation. Then, one logical connective after another is eliminated upwards along the chain, as is shown in below examples. If the proof is successful, the sequents at the top of the proof are all identify axioms in the form of (15).

(15) Axiom: $A \vdash A$

By convention, $p, q, r \dots$ represent atomic propositional letters, A, B, C represent any (propositional) formulas, and X, Y, Z represent sets of such formulas. I omit the set notations both in the antecedent (i.e. the left-hand) side of each turnstile and the succedent (i.e. the right-hand) side. (16) shows the axioms for the connectives, & and $\rightarrow$ (I use $\wedge$ for & in Gentzen sequent presentation for some technical reasons). I omit the rules for other connectives. *Cut* is an admissible rule⁴³ which is not necessary for the proof system, but is useful for improving the efficiency of the proof.

(16) Logical rules:

$$\begin{array}{ll}
 \text{a. } \frac{A, B \vdash X}{A \wedge B \vdash X} \wedge L & \frac{X \vdash A \quad Y \vdash B}{X, Y \vdash A \wedge B} \wedge R \\
 \text{b. } \frac{X \vdash A \quad Y, B \vdash Z}{X, Y, A \rightarrow B \vdash Z} \rightarrow L & \frac{X, A, Y \vdash B}{X, Y \vdash A \rightarrow B} \rightarrow R \\
 \text{c. } & \frac{X \vdash A \quad A \vdash Z}{X \vdash Z} \text{Cut}
 \end{array}$$

I have omitted some of the ‘contextual’ structural variables (i.e. $X, Y, \dots$) for readability. Note that $\wedge L$ in (16a) is ‘pure’ in the sense that it does not introduce a new propositional variable in the inference from the top to the bottom, as opposed

⁴³ That is, if any sequent that we can prove with Cut is provable without Cut.

to the ‘impure’ inference rule from $A \vdash X$ to $A \wedge B \vdash X$, which is valid in classical logic, but has incorporated structural weakening in the Antecedent. The ‘pure’ presentation is preferable for the reason that we have discussed in section 4. In (16), except for the Cut rule,⁴⁴ the number of the connectives decreases by one along each consecutive step upwards. Because there are only a finite number of connectives in each sequent to be proved, any proof is decidable in a finite step, unless Cut is used.

Remember the successive use of $\&I$ ($= \wedge R$ here) and $\&E$ ($= \wedge L$), which may allegedly lead to an infinite inference. With Cut, this claim is substantiated, as in (17).

(17) *Proof 1*

$$\frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash p \wedge q} \wedge R \quad \frac{\frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash p \wedge q} \wedge R \quad \frac{p \wedge q \vdash p \wedge q}{(p \wedge q), (p \wedge q) \rightarrow r \vdash r} \wedge L \quad r \vdash r}{(p \wedge q), (p \wedge q) \rightarrow r \vdash r} \rightarrow L}{p, q, (p \wedge q) \rightarrow r \vdash r} Cut$$

(18), in which Γ and Δ represent the two sub-proofs of (17), represents the proof in (17) in brief. If the Cut rule is used, then this proof might not terminate in a finite step, given the sequent to prove, $p, q, (p \wedge q) \rightarrow r \vdash r$.

(18) *Proof 1 (with abbreviation)*

$$\frac{\Gamma \quad \Delta}{p, q, (p \wedge q) \rightarrow r \vdash r} Cut$$

In the position of the sub-proof Γ in proof 1, one could insert a larger sub-proof, e.g., the whole of the proof 2 in (19).

⁴⁴ Cut is not a logical rule (which is a rule for a truth functional connective/operator). Its inclusion in (16) is for presentational convenience only.

(19) (Sub)-proof 2

$$\frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash p \wedge q} \wedge R \quad \frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash p \wedge q} \wedge R \quad \frac{p \wedge q \vdash p \wedge q}{p, q \vdash p \wedge q} \wedge L}{p, q \vdash p \wedge q} Cut$$

Note that the left premise and the conclusion of Cut are both $p, q \vdash p \wedge q$. Thus, we can use the conclusion sequent as a left premise of another Cut, by repeating the whole of the right premise of the original Cut as the right premise of this additional Cut. Thus, there is no maximal limit to the size of the sub-proof in (19), leading to the infinity (or undecidability) problem.

However, as Girard (1987) and others showed, Cut is an admissible rule in Gentzen sequent presentation. Without Cut, Proof 1 is represented as Proof 3.

(20) Proof 3 (Without Cut)

$$\frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash p \wedge q} \wedge R \quad r \vdash r}{p, q, (p \wedge q) \rightarrow r \vdash r} \rightarrow L$$

Other than Cut, all the rules in (15)~(16) reduce the number of connectives by one along each consecutive step upwards, and thus, all the proofs are decidable in finite steps. Consequently, successive use of $\wedge L$ and $\wedge R$ does not lead to an infinite inference.

Finally, (22) show that the equivalence in (8), repeated here as (22), is provable only with $\&I$ (or $\wedge R$ here) as a rule of the logic. The proof in (23a) requires $\wedge R$ in the top left sub-proof.

(22) $(p \wedge q) \rightarrow r \vdash p \rightarrow (q \rightarrow r)$ (23) a. $\vdash$

$$\frac{\frac{\frac{p \vdash p \quad q \vdash q}{p, q \vdash (p \wedge q)} \wedge R \quad r \vdash r}{p, q, (p \wedge q) \rightarrow r \vdash r} \rightarrow L \quad \frac{p, (p \wedge q) \rightarrow r \vdash q \rightarrow r}{p, (p \wedge q) \rightarrow r \vdash p \rightarrow (q \rightarrow r)} \rightarrow R}{p, (p \wedge q) \rightarrow r \vdash p \rightarrow (q \rightarrow r)} \rightarrow R$$

$$\begin{array}{c}
\text{b. } \neg \\
\frac{p \vdash p \quad \frac{q \vdash q \quad r \vdash r}{q, q \rightarrow r \vdash r} \rightarrow L}{p, q, p \rightarrow (q \rightarrow r) \vdash r} \rightarrow L \\
\frac{p \wedge q, p \rightarrow (q \rightarrow r) \vdash r}{p \wedge q, p \rightarrow (q \rightarrow r) \vdash r} \wedge L \\
\frac{p \wedge q, p \rightarrow (q \rightarrow r) \vdash r}{p \rightarrow (q \rightarrow r) \vdash (p \wedge q) \rightarrow r} \rightarrow R
\end{array}$$

In this section, I showed that successive use of &I(=∧R) and &E(=∧L) does not lead to an infinite inference in Gentzen sequent presentation without Cut. I also showed that we need &I as an inference rule to support CMPP as an application rule in spontaneous inference. I did not show how we can prevent infinity which could be induced by the use of structural rules (such as expansion and weakening) in the proof presentations, but for some rough ideas (in the context of Modal logic), see Hudelmaier (1996).

6 Loose ends and speculations

This section deals with some loose ends. The discussion will be mostly speculative and incomplete.

6.1 Mental logic over non-propositional representations.

As I wrote in section 3 (cf. footnote 23), Sperber & Wilson make their mental logic operate over ‘non-propositional representations,’ and I add some comments to this claim.

One can interpret this claim in two different ways. One interpretation is that, in Sperber & Wilson’s inference system, propositional letters are not always interpreted in terms of their truth values. With this interpretation, it is misleading to state that their mental logic operates over ‘non-propositional representations,’ because the underlying system may still be propositional logic, only with different resource management properties (or, less technically, with different ‘interpretations’ of propositional formulas). Thus, we can still stay inside propositional logic, only with varying ways of interpreting the propositional language.

Though truth based inferences may not be the only kind of inferences for Relevance Theory, they are still an important target of their pragmatic analysis. Thus, we can safely assume that part of the tasks of their mental logic is to explain one’s spontaneous truth-based inferences. But then the incompleteness problem, among others, is as serious a problem in their system as in an inference system with

the truth-based semantics as the ‘only’ intended interpretation. Because of this, the claim that their mental logic operates over ‘propositional’ letters that may be interpreted in a different way from their truth values does not make their system without &I any less problematic. It will still underachieve its intended tasks (or they may need to add stipulative side conditions to make it work in a complete way).

On the other hand, it is true that such a claim makes the application of my proposal less straightforward. Instantiation of the proposal in S&W’s inference system requires further research. In this section, I only sketch a speculative way of applying the proposal to such a multi-purpose inference system.

As long as S&W can use the same inference language, only interpreted in varying ways depending on which kinds of inferences they are dealing with, we can keep the logical language more or less the same as the one in classical propositional logic. To simplify things, let me stick to the propositional logic language that we have used in this paper, made out of expressions such as P , $P \& Q$, $P \rightarrow Q$ etc. To instantiate a multi-purpose inference system as above, we may see the expressions in the logical language in multi-modal interpretations.⁴⁵ That is, in one mode of interpretation, one will interpret the formulas in terms of their truth values (then in this mode of interpretation, the inference system is basically the classical propositional logic). In another mode, one may interpret P , Q etc. as ‘resources,’ as in linear logic (then in this mode, one occurrence and two occurrences of the same formula, say, P , make a difference, just as one bottle of beer and two bottles of beer as resources are interpreted differently). With such multi-level interpretations, as long as the inference system is equipped with the whole set of introduction and elimination rules for all the connectives in the mode of truth based interpretation, one can at least confirm that the system does its job in a complete way as far as the truth based inference is concerned. But we can also propose a multi-modal system in which the system is sound and complete in every mode, with regard to the intended interpretation in each mode (that is, the system will do its job in a complete way in each mode of interpretation). As I sketched in section 4, the alleged infinite inference problem is caused by simplistic application of structural rules. Thus, the multi-modal system may solve this problem without stipulatively banning introduction rules for truth functional connectives. Instead, we can choose the right set of modes of interpretations with appropriate structure management properties as its sub-systems.

⁴⁵ To have the interpretations in different modes interact with one another, we would have to modify the (logical) language expressions in the system, such as introducing some ways of signifying the different modes of interpretations or defining interaction rules between different modes, where such interaction rules might in turn require an addition of modal operators into the logical language. I leave a more accurate exposition of such a multi-modal system for another occasion.

As I said above, I only provide a sketch of how to instantiate this paper's proposal in such a multi-modal system. First, with regard to the 'fully propositional status' of the linguistic meaning of some language expression, we may simply base such status on the truth-evaluability of the expression in its interpretation in the truth-based mode of interpretation.⁴⁶ That is, the rules for propositional connectives/operators cannot apply in any mode until all the premises of the rules are judged to be fully truth evaluable in the truth based mode of interpretation. This will solve the alleged overgeneration in enrichment case. What is more difficult is how to control the alleged infinity in terms of structural rule application (again, I assume that the problem is not caused by introduction rules per se). This will require more careful work, because, by definition, structural properties of logical expressions vary across various modes of interpretations. However, the different structure management properties in different modes mean that we can naturally get rid of certain structural rules in certain modes of interpretations without stipulatively banning those rules (e.g. in Intuitionistic logic, weakening in the succedent is impossible, and thus, the alleged infinity in (11b) simply does not arise).

In this paper, I do not investigate whether such mode internal variation of structural properties is enough to cover all the kinds of inferences that Sperber & Wilson aim to explain. For example, an interesting question with regard to the multi-level interpretations of the logical expressions is whether we need to include a mode of interpretation in which the expressions are interpreted as tokens. But note that even in this mode, it does not make sense to see '&' as a token as well, if we still see it as a 'deductive' system.⁴⁷ Remember that formulas such as p, q, r on the one hand, and the connectives, operators such as $\&, \vee$ have different statuses in logic. For example, only the former can stand on their own as well-formed expressions in the language, whereas the latter could not do this, as shown by the

⁴⁶ I abstract away from the mapping between natural language expressions used in utterances and the corresponding expressions in logical languages because it requires a more complex exposition. Roughly, the statement in the main text would change into the following: "we may simply base the truth evaluability of a language expression uttered (i.e. a 'sentence' in the case of a trivial proposition) on whether the corresponding logical expression (i.e. a formula in the case of a trivial proposition) can play a non-trivial role in the truth based inference in the context of the utterance." To remind the reader of my argument in section 3, one can decide whether the logical expression in question can play a non-trivial role or not in such inferences without letting it interact with other contextually available propositional formulas in terms of truth based inference rules. See section 3 for details.

⁴⁷ We could stop seeing this level or mode as part of the deductive system, but then it would become impossible to relate this mode to the other modes which are defined to be deductive. As an example, note that in Lambek calculus, which pairs LF as a relational structure with PF as a relational structure, phonological strings include (interpretations of) 'logical connectives,' such as the binary connective $\cdot$ which connects, *the* and *boy* producing, (*the* $\cdot$ *boy*).

ill-formed expressions, $\&$, $\vee p$ and $q \rightarrow$. These connectives are not independent elements in the language; they play a role of mapping (simpler) formulas to (more complex) formulas in terms of functional derivability, where the derivability is inherently related to the basic property of the deductive system. Thus, insisting on the difference between p , q on the one hand, and $(p\&q)$ (in the Antecedent) just because only the latter has $\&$ as a token misses the point. If we interpreted $\&$ etc. as a token (and if we did not add proper logical connectives instead of them, which would support the 'token' based deductive system), then the system would simply stop being deductive. The interpretations of the connectives are inherently related to the properties of the deductive system that makes use of them. That is why 'commas' in the Antecedent of a sequent have the same interpretation as $\&$ and 'commas' in the succedent of a sequent have the same interpretation as $\vee$. Just like 'commas' in proof representations cannot be treated as tokens, the connectives whose interpretations are inherently related to those commas cannot be treated as tokens. From a different viewpoint, propositions or any well-formed formulas can be interpreted as tokens because they are part of the set of well formed expressions in the language. On the other hand, the connectives such as $\wedge$, $\vee$ and $\rightarrow$ are not well-formed expressions on their own. They are defined to express the derivability relations supported in the chosen deductive system. Thus, if they were treated as tokens and were not assigned the semantics that are expected from the basic properties of the deductive system, then the system would stop becoming deductive. This means that even in the interpretation of logical expressions as tokens, we cannot treat $\&$ as a token. We would have to define some functional interpretation which maps the token interpretations of p and q to the token interpretation of $p\&q$.⁴⁸ I leave the structural property in that mode of interpretation for future research.

The second way of interpreting the claim that Mental Logic operates over non-propositional expressions is that their inference system deals with sub-propositional expressions (such as individual terms and predicate expressions), as well as propositional expressions. But we can accommodate this requirement just by using a system such as predicate logic (or a variant of higher order logic) at sub-propositional levels. In this paper, I did not explicitly look into substructures of the propositional expressions such as P , Q , R , but we presupposed the necessity of sub-propositional expressions in the inference language in many parts of the paper. For example, in order to model entailment relations at sub-propositional level in terms of set containment relations, I would need to use sub-propositional expressions in the logical language. Incompleteness at the sub-propositional level is a separate

⁴⁸ As in linear logic, we would need different kinds of $\&$, such as a multiplicative one as opposed to an additive one.

issue that I abstract away from.⁴⁹ At the moment, as long as the inference system is complete at the level of the truth-based propositional calculus, it serves our purpose (note that in the standard predicate calculus, the connectives $\&$, $\vee$, $\rightarrow$ etc. are still interpreted as propositional/truth functional connectives. See footnote 31).

⁴⁹ This does not mean that incompleteness does not matter at the sub-propositional level. Rather, it is not yet clear to me how we can use deductive inferences to explain spontaneous inferences at that level. A Lambek Calculus NL, which regards syntactic categories as formulas, such as $\text{NP} \rightarrow \text{S}$, S , $\text{NP} \bullet \text{NP}$ (cf. $p \rightarrow q$, q , $p \& q$, respectively), has been shown to be complete with regard to a well-defined semantic model (i.e. free groupoid), and this shows that it is possible to achieve completeness with regard to some well-defined semantics by providing formulas-as-types to both propositional and sub-propositional logical expressions. However, in our case, the inference is not about compositional derivations of LF representations from the lexical level (which the use of logic in the syntax can take care of). Rather, it is about spontaneous inferences manipulating the out-put of the syntactic derivation at LF after its near isomorphic translation to Language of Thought (LoT) representations. I am not sure how the use of deductive systems can constrain spontaneous inferences at the sub-propositional level in LoT. Lexical enrichment is an example of such sub-propositional inferences. However, whereas ‘narrowing’ in lexical enrichment might be captured in terms of set-containment relation between sub-propositional concepts, it is not clear how using predicate logic (or higher order logic) can provide insight in the process of loosening. If we use first-order predicate logic or a variant of simply typed lambda expressions to represent LoT, then to the degree that the logical language expressions are constrained by some formal properties of the language, spontaneous inferences using such language expressions will be constrained as well. However, as one can informally understand by considering how the use of English may constrain one’s general thinking in English, the way that one’s inference is constrained by the language that one uses in inference (because of the limited expressive power of the language) may be quite different from the way that the classical propositional logic may constrain one’s spontaneous propositional inferences as has been discussed in this paper. One possibility is that ‘sub-propositional inferential processes’ involved in enrichment etc. are not deductive in a direct way. That is, to the degree that both the starting point and the end result of (lexical) enrichment are part of propositional representations which have some formal structures as are expressible in simply-typed lambda expressions, for example, enrichment may still be constrained by the syntax of such logical expressions, but again, that is a different kind of constraints from the constraints that (the deductive rules of) the classical propositional logic may provide for propositional inferences. In this interpretation, the process of enrichment is not explainable in terms of a logical rule such as $\&\text{I}$, MPP etc. It is only that enrichment manipulates some logical language representations. We can explain possible use of set-containment relations in enrichment (i.e. narrowing) in this indirect way. That is, enrichment may manipulate certain (logical) properties of the language representations that it operates over, and also, because the final product of sub-propositional inferences is the proposition expressed which does enter into properly deductive propositional inferences, enrichment may be geared towards the preservation of ‘logical entailment relation’ (again, mimicked by set-containment relations) at the sub-propositional level. However, the process of enrichment itself might still not be definable as a deductive rule. I leave this issue for future research.

6.2 Denotational or procedural views of semantics and soundness and completeness of the system

In this paper, I adopted a ‘denotational view’ of semantics. That is, our semantics modelled the denotations of the logical expressions, rather than modelling the syntactic proofs/derivations themselves, as Heyting did (see the next paragraph). Braine and O’Brien (e.g. O’Brien 2004), on the other hand, explicitly advocate a kind of ‘procedural semantics.’ The claim is that the semantics of the connectives are based on what they allow one to do with them.

There are two ways of interpreting this claim. In one interpretation, their semantics is based on what mental logic rules allow one to do with them *in the semantics*. In this case, however, there is no inherent difference between their conceptions of the semantics and the above mentioned ‘denotational’ view of the semantics. In the denotational view, the semantics of logical language models what the syntactic system can do by way of interpreting the syntactic objects in the intended semantic structure. In that sense, the semantics of the deductive system does correspond to what the syntactic system allows one to do in the semantics. Whether or not they use truth tables in the intended semantics is a separate issue. We could interpret the inference language in a semantic structure that is different from the one represented by truth tables (i.e. the Boolean lattice). Braine and O’Brien could define whatever semantic structure is suitable for their purpose as long as the intended semantics is formally well-defined. But whatever denotations they may assign to the inference language, they must check whether what the system allows one to do at the level of syntactic derivations matches up with what the system allows one to do at the level of the (system internal) denotations in a sound and complete way, so that one can confirm that the system actually does the job that they intend it to do.

The alternative interpretation of their claim is more interesting. They might be assuming that their semantics directly model their ‘proofs’ (or their deductive steps). This reminds me of Heyting’s semantics. In interpreting propositional languages, Heyting did not try to find out when each propositional formula is true. Instead, he tried to find out what the proof of each formula is (cf. Girard 1989: 5). Thus, Heyting first stipulated that the interpretation of each atomic formula (say, P) is its proof.⁵⁰ After that, he stipulated that a proof of $P \wedge Q$ is a pair (p, q) consisting of a proof p of P and a proof q of Q (cf. Girard, 1989:5).

⁵⁰ What counts as a proof of an atomic formula is not clear, but consider what one would do to prove each P , such as $2+3=5$. Probably one could place two objects on the table, add three more objects to them, and then count the total as five, which may count as a proof of $2+3=5$. In any case, the point of the direct interpretation of proofs is that, once we agree upon the proof of each atomic formula as an interpretation of that formula, then we can compute the proofs of more complex formulas out of the proofs of atomic formulas at the level of model structure, in the way

However, I do not think that an attempt to interpret the syntax of their mental logic in this direct way would be successful for Braine and O'Brien for various reasons. First, in the case of Heyting, he was interested in the direct interpretations of the proofs themselves. Thus, for Heyting, it does not matter what the resultant semantics turns out to be, as long as it directly represents the syntactic proofs/derivations (and as long as the semantic structure turns out to be well-defined). This is not the case for Braine and O'Brien, they have some empirical phenomena to explain, that is, spontaneous inferences as psychological phenomena. Thus, their semantics should have an appropriate structure as a model of one's system for spontaneous inferences.⁵¹ Secondly, the incompleteness of the system without &I⁵² comes from the 'incompleteness' of the syntactic system at the basic level.⁵³ Given this inherent incompleteness of their syntactic system, direct interpretation of their system is not likely to help Braine & O'Brien's system (which does not have the property of symmetry) with regard to soundness and completeness relative to the direct interpretation. This is because the semantic structure is evaluated with full generality. Even though the initial interpretation rules map only the objects that their syntactic rules allow one to generate onto some

suggested in the main text. Note that even in this direct interpretation method, Heyting had to provide some abstract 'denotations' to the language expressions, whether they are atomic or complex. The difference between the 'denotational view' and the direct/syntactic view of semantics then is simply that for the denotational view, one provides a matching semantic structure to a syntactic system (such as a Boolean lattice to the classical propositional logic) in the first place and then tries to prove soundness and completeness of the syntactic system with regard to that semantics, or for the syntactic view, one tries to directly represent each syntactic object and all the proof steps in one's semantic model, hoping that the resultant semantics constitutes a well-defined structure. The benefit of the first strategy is that the semantic structure is already well-defined at the start, because one picks up a well-defined structure in the first place. But soundness and completeness might not hold (and then one might look for another well-defined semantic structure as a candidate). With the latter view, one tries to set up the semantics in such a way that it follows each syntactic proof step. Ideally, the syntax should become sound and complete with regard to the resultant semantics created in this way. However, a problem of this second strategy is that there is no confirmation that one can create a well-defined semantic structure at the end of the day. Also, for technical reasons, maintaining soundness and completeness turns out not to be so easy to sustain even in this way of setting up the semantics tailor-made for the syntax. See chapter 1 of Girard (1989) in this regard.

⁵¹ In other words, though the (system-internal) semantics of a logical language is independent of the (system-external) semantics as is represented in the inference data, these two kinds of semantics should match up quite closely (ideally, they should be isomorphic to one another) so that one can use the system in an empirically meaningful way.

⁵² Or in Braine and O'Brien's theory, the incompleteness of the system which puts &I into a different group from the core group of rules. See the next subsection, 6.3.

⁵³ Informally, a syntactic system itself would become 'incomplete' if it is equipped with only one of the pair of rules for a connective used in the language, unless this elimination falls out from the basic structural property of that language.

semantic objects (and thus, their syntax will be sound and complete with regard to the semantics at this initial stage), evaluation of the resultant semantic structure with its fully general representational capacity will justify addition of some more semantic objects whose syntactic correspondents the syntax cannot generate with the given set of rules.

I wait for another occasion to provide a full review of Braine and O'Brien's analysis. In this section, I have added some speculative comments about alternative ways of interpreting deductive systems.

6.3 Multi-modal inference system.

As I mentioned in section 4 (footnote 35), Braine and O'Brien divide mental logic schemas into different categories. The basic ones (such as MPP) apply automatically. Some others (including &I) are only applied if their output will feed one of the basic ones. They claim that this solves the alleged infinity problem. For example, one cannot apply &I as in $P \vdash P \& P$ unless this application feeds one of the main schemas, such as MPP, avoiding the alleged infinity in (11a) in section 4.⁵⁴ I do not review this proposal in detail in this paper, but there are several problems. First, as I show in the main text, the alleged infinity, even if it exists, is not caused by introduction rules themselves. Thus controlling the use of introduction rules does not solve the problem in a complete way (without further stipulations that make the system even more complex). Also, dividing the rules for logical connectives into different groups is dangerous, because there are certain derivability relations among these rules and separating them into groups with restricted feeding relations risks making the system incomplete not only in each group but at the level of the whole system.⁵⁵ Compare their proposal with the informal suggestion in section 4 in which one controls structural rule application. To require that one may apply structural rule only if the output feeds into a logical rule is less harmful in several ways. First, the division between logical rules (i.e. rules for logical operators connectives, such as &, $\vee$, and $\rightarrow$) and structural rules (i.e. weakening, contraction, etc.) are already there in logic. There are several

⁵⁴ Without further restrictions, Braine and O'Brien's system does not solve the alleged overgeneration with enrichment as in (2)~(4) because in that case, the output of &I does feed MPP. But use of &I in relevance theory has not been one of their concerns.

⁵⁵ For each connective, having the elimination rule without the introduction rule is problematic, as I have shown in the paper. This is the same (in a slightly different way) if the introduction rule and the elimination rule are put into different groups of rules. Also, consider the equivalence between $(P \rightarrow Q)$ and $(\neg P \vee Q)$. Given equivalence relations like this one, putting the rules for $\rightarrow$ in one group, and the rules for $\vee$ in another, restricting the use of the latter, then also risks making the system incomplete with regard to the intended interpretation.

diagnostics that one may use for telling the differences between them. For example, simply count the number of connectives before and after each rule application. Logical rule application necessarily influences the number of connectives (or, equivalently, it influences the complexity of the formulas or the structured configurations of formulas). On the other hand, application of a structural rule in itself does not influence the complexity of the formulas.⁵⁶ It is not only that logical rules and structural rules are different in nature. See the footnote 36 for an observation to the effect that logical rules are independent of the structural properties of the system. This independence of logical rules from structural properties allows us to limit the application of structural rules in the way that we explained in section 4. In fact, there are several established proofs that show that applying a structural rule only if the output feeds into some logical rule does not influence the set of derivable (or provable) sequents (see Hudelmaier 1996, for example, in this regard). Thus, we have some confirmation that restricting the application of structural rules in this way to avoid the alleged overgeneration does not influence the derivability of sequents (or, semantically, the validation of arguments) in the inference system. If the original system without such control of structural rule application is complete with regard to the intended interpretation, then the same system with such control is also complete.

I leave further evaluation of Braine and O'Brien to further research.

6.4 Truth-based judgment in spontaneous inference

In (7) (repeated here as (24) below) in section 3, I argued that the spontaneous system at the basic level should be equipped with the rule which can directly support the truth functionally equivalent role that p, q on the one hand, and $(p \& q)$ on the other, play in the Antecedent of a sequent.

- (24) a. $p, q, (p \& q) \rightarrow r \vdash r$
 b. $p \& q, (p \& q) \rightarrow r \vdash r$

Carson (p.c.) claims that it is not clear why this equivalence is something that the spontaneous inference system should be expected to explain. I agree that it is

⁵⁶ Contraction of $P \& Q, P \& Q \vdash X$ to $P \& Q \vdash X$ may seem to influence the number of connectives, but this reduction of the complexity of the form is not really because of the structural rule application itself. Note that from $P \& Q, P \& Q \vdash X$, one may first eliminate two occurrences of $\&$ via $\&E$, producing, $P, P, Q, Q \vdash X$. After that, one can apply contraction, producing, $P, Q \vdash X$. Based on the assumption that different routes to reach the proof of the same sequent actually represents the same proof, we can claim that structural rules do not influence the complexity of the structured configurations of formulas

debatable whether we can directly recognize the same semantic roles that p , q separately and $p \& q$ together play as premises in our spontaneous truth-based inference at the data level. However, I am not discussing the thing with regard to the semantics of the inference data only. I am also evaluating the spontaneous inference system with regard to whether it does its job in a complete way as a well-defined deductive system. A system that lacks $\&I$ but with $\&E$ cannot do $\&I$ in the syntax by stipulation, but the intended semantics (if the Boolean semantics as I suggested is the intended semantics) predicts that the system can validate that syntactically impossible sequent in the semantics. Thus, the system is inconsistent between the verdict in the syntax and the (contradicting) verdict in the intended semantics. They could provide an alternative semantics as the intended semantics so that this inconsistency can be resolved, but for the moment, I find it hard to come up with such an alternative which matches with their syntax in a complete way. Also, even if they could successfully provide the semantics with regard to which their suppression of $\&I$ from the inference system is complete, that alternative would mean that we could only recognize the above mentioned equivalence in the roles played by p , q and $p \& q$ only in an indirect (or reflective) way as I explained in section 3. I am not sure if I feel as if my own recognition is only indirect in my spontaneous inference. Given that it makes the definition of the intended semantics far more difficult, I argue that it is better to equip the inference system with $\&I$ at the base level, and explain why $\&I$ is not used in our spontaneous inference in certain cases for independent reasons. For example, as I have sketched in 6.1 and 6.3, with multi-modal interpretations of the inference language, we may recognize the difference between p , q and $p \& q$ as premises in a mode of interpretation that is different from the truth based one. This may explain why we feel as if there are differences between these two at the level of intuitions.

7 Conclusion

If a pragmatic inference system is to explain one's truth-based inference (possibly among other kinds of inference), it is not desirable to eliminate logical introduction rules completely from the inference system, with view to preserving the consistency of the system as a whole. Use of introduction rules in the inference system as a whole does not lead to overgeneration via enrichment. Introduction rules can apply only with fully propositional elements as premises, and thus, such rules cannot be applied before the recovery of the proposition expressed. The alleged infinite inference steps are not caused by introduction rules per se, and the problem must be solved independently.

References

- Avron, Amon. (1993). Gentzen-Type Systems, Resolution and Tableaux. *Journal of Automated Reasoning*, 10, 265-281.
- Belnap, Nuel. (1996). The Display Problem. In Heinrich Wansing (ed.), *Proof Theory of Modal Logic*. Dordrecht: Kluwer, 79-92.
- Carston, Robyn. (2002). *Thoughts and Utterances: the pragmatics of explicit communication*. Oxford: Blackwell.
- Došen, Kosta (1988). Sequent systems and groupoid models I. *Studia Logica*, 47, 353-389.
- Girard, Jean-Yves. (1987). *Proof theory and logical complexity*. Amsterdam: Elsevier.
- Girard, Jean-Yves. (1989). *Proofs and Types*, Translated and with appendices by Paul Taylor and Yves Lafont. Cambridge: Cambridge University Press.
- Hall, Alison. (2006). Free enrichment or hidden indexicals. *UCL Working Papers in Linguistics* 18:71-102.
- Hudelmaier, Jörg. (1996). A contraction-free sequent calculus for S4. In Heinrich Wansing (ed.), *Proof Theory of Modal Logic*. Dordrecht: Kluwer, 3-15.
- Johnson-Laird, Philip N. (1997). 'Rules and illusions: a critical study of Rips' "the psychology of proof."' *Minds and Machines* 7.3:387-407.
- O'Brien, D.P. (2004). 'Mental-logic theory: What it proposes, and reasons to take this proposal seriously.' In J. P. Leighton & R.J. Sternberg (Eds.), *The nature of reasoning*, Cambridge: Cambridge University Press, 205-233.
- Partee, Barbara and Mats Rooth. (1983). 'Generalized conjunction and type ambiguity.' In Rainer Bäuerle, Christoph Schwarze, and Arnim von Stechow (eds.), *Meaning, Use, and Interpretation of Language*. Berlin: de Gruyter, 361-393.
- Rips, Lance. J. (1997). 'Goals for a theory of deduction: reply to Johnson-laird.' *Minds and Machines* 7.3:409-424.
- Sperber, Dan. and Deirdre Wilson. (1986/1995) *Relevance: Communication and Cognition*. Oxford: Blackwell.
- Stanley, Jason. (2000). Context and logical form. *Linguistics and Philosophy* 23:391-434.
- Stanley, Jason. (2002). Making it articulated. *Mind and Language* 17 1& 2:149-168.
- Takeuti, Gaisi. (1987). *Proof Theory: Studies in logic and the foundations of mathematics, Vol 81*. North-Holland: Elsevier Science Ltd; 2nd edition.
- Wansing, Heinrich. (1993). *The Logic of Information Structures*. Lecture Notes in AI 681, Berlin: Springer-Verlag.
- Wansing, Heinrich. (1998). *Displaying Modal Logic*, Dordrecht: Kluwer.