



Fundamentals of the CTT Approach

Shahid Rahman, Nicolas Clerbout

► To cite this version:

Shahid Rahman, Nicolas Clerbout. Fundamentals of the CTT Approach. Master. France. 2015.
cel-01228820v3

HAL Id: cel-01228820

<https://shs.hal.science/cel-01228820v3>

Submitted on 23 Nov 2015

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

**THE FOLLOWING TEXT IS AN OVERVIEW ON EXISTING LITERATURE:
THERE IS NO CLAIM OF ORIGINALITY**

Fundamentals of the CTT Approach¹

By Shahid Rahman* and Nicolas Clerbout**

(*Univ. Lille, CNRS, UMR 8163 - STL - Savoirs Textes Langage, F-59000 MESH, Lille, France²,
Universidad de Valparaíso)

(shahid.rahman@univ-lille3.fr)

Within Per Martin-Löf's Constructive Type Theory (CTT for short) the logical constants are interpreted through the Curry-Howard correspondence between propositions and sets. A proposition is interpreted as a set whose elements represent the proofs of the proposition. It is also possible to view a set as a problem description in a way similar to Kolmogorov's explanation of the intuitionistic propositional calculus. In particular, a set can be seen as a specification of a programming problem: the elements of the set are then the programs that satisfy the specification (Martin-Löf 1984, p. 7). Furthermore in CTT sets are also understood as types so that propositions can be seen as data (or proof)-types.³ We will start by a quick over-view of the principles of the CTT approach. We will then recall the rules of intuitionistic predicate logic in CTT.

The general philosophical idea is linked to what has been called the fully interpreted⁴ approach in which special care is taken, in Martin-Löf's own words (1984, p.2), to

avoid keeping content and form apart. Instead we will at the same time display certain forms of judgement and inference that are used in mathematical proofs and explain them semantically. Thus, we make explicit what is usually implicitly taken for granted”.

The *explicitation* task involves bringing into the object language level features that determine meaning and that are usually formulated at the metalevel.

According to the CTT view, premises and conclusion of a logical inference are not propositions but judgments.

A rule of inference is justified by explaining the conclusion on the assumption that the premises are known.

Hence, before a rule of inference can be justified, it must be explained what it is that we must know in order to have the right to make a judgment of any one of the various forms that the premises and conclusion can have (Martin-Löf 1984, p.2).

1 Translated from the Spanish by the Hanna Karpenko from the chapter on an overview of CTT in the book *Las Raíces Dialógicas de la Teoría Constructiva de Tipos*, by Rahman/Clerbout, London : College Publications, in preparation.

² The paper is a draft for first contribution the ANR research program SÊMAINÔ, and is being developed in the context of the axe transversal *Argumentation* of STL and the research axe ADA of the MESHs.

³ Cf. Nordström et al. (1990) and Granström (2011).

⁴ For a thorough discussion on this issue, see Sundholm (1997, 2001).

Two further basic tenets of CTT are the following:

1. No entity without type
2. No type without identity

Accordingly, we can take the assertion that an individual is an element of the set A as the assertion that that individual instantiates or exemplifies type A . But what is a type, and how do we differentiate between its examples and those objects that are not of a given type? Or, more fundamentally, what must we know in order to have the right to judge something as a type?

In CTT, a set is defined by specifying its *canonical elements*, the elements that “directly” exemplify the type, and its non-canonical ones, the elements that can be shown, using some prescribed method of transformation, that they are equal (in the type) to a canonical one. Moreover, it is required that the equality between objects of a type be an equivalence relation.⁵ This is what the second tenet is about: the introduction of an equivalence relation in a set (an object of the type ‘set’).

When we have a type we know, from the meaning explanation of what it means to be a type, what the conditions are for being an object of that type. So if A is a type and we have an object b that satisfies the corresponding conditions, then b is an object of type A , which we formally write $b : A$.⁶ In such a frame

$b : A$

$A \text{ true}$

can be read as

b is an element of the set A

A has an element

b is a proof of the proposition A

A is true

b fulfills the expectation A

A is fulfilled

b is a solution to the problem A

A has a solution

‘Set’ itself does not instantiate the type *set*, since we do not have a general method for generating all possible ways of building a set. However, given the type ‘set’ we can build the objects that instantiate it by the means described above.

The four basic forms of judgements in CTT are

$A : \text{set}$

$A = B : \text{set}$

$A : \text{prop}$

$A = B : \text{prop}$

In higher type hierarchy we also have the judgements of the following form:

$\alpha : \text{type}$

⁵ For a thorough discussion see Granström (2011, pp. 54–76).

⁶ Martin-Löf used the sign “ $\in$ ” in order to indicate that something, say a , is of some type, say B . He even suggests to understand it as the copula ‘is’. Nordström et al. (1990) also uses this notation while other authors such as Ranta (1994) use the colon. Granström (2011) distinguishes the colon from the epsilon, where the first applies to non-canonical elements and the latter to canonical ones. We will use the colon.

$$\alpha = \beta : \text{type}$$

Hypotheticals: The judgements we have introduced so far do not depend on any assumptions. They are *categorical* judgements. The CTT language has also *hypothetical* judgements of the form:

$$B \text{ type } (x : A)$$

Where A is a type which does not depend on any assumptions and B is a type when $x : A$ (the *hypothesis* for B). In the case of sets we have that b is an element of the set B , under the assumption that x is an element of the A :

$$b : B (x : A) \quad (\text{more precisely: } b : el (B) (x : el (A)))$$

The explicit introduction of hypotheticals carries with it the explicit introduction of appropriate substitution rules, but we will deal with them further on in our text, when we present the CTT view on equality.

As pointed out by Granström (2011, p. 112) the form of assertion $b : B (x : A) \quad (b : B (x : A))$ can be generalized in three directions:

1. Any number of assumptions will be allowed, not just one;
2. The set over which a variable ranges may depend on previously introduced variables;
3. The set B may depend on all introduced variables

Such a list of assumptions will be called a *context*. Thus we might need the forms of assertion

$$b : B (\Gamma) - \text{where } \Gamma \text{ is a context (i.e., a list of assumptions)} \\ \Gamma : \text{context}$$

What is *given in a context*, the *given contextually-dependent knowledge*, is whatever can be inferred from the hypotheses constituting the context. In fact, it is usually distinguished between what is *actually given* in the context (*actual knowledge*), namely the variables themselves and the judgements involving these variables and what is *potentially given* (*potential knowledge*), namely what can be inferred by the rules of type theory from what is actually given.

Assume $A_1 : \text{set}$ and $A_2 : \text{set } (x : A_1)$.

The judgement

$$A_3 : \text{set } (x_1 : A_1; x_2 : A_2)$$

means that

$$A_3 [a_1/x_1; a_2/x_2] : \text{set whenever } a_1 : A_1 \text{ and } a_2 : A_2 [a_1/x_1]. \\ A_3 [a_1/x_1; a_2/x_2] = A_3 [a'_1/x_1; a'_2/x_2] : \text{set} \\ \text{whenever} \\ a_1 = a'_1 : A_1 \text{ and } a_2 = a'_2 : A_2 [a_1/x_1].$$

We say that A_3 is a family of sets over A_1 and A_2

Hence, a hypothetical judgement is made in a context that has the form:

$$x_1 : A_1, x_2 : A_2, \dots x_n : A_n$$

where we already know that A_1 is a type, A_2 is a type in the context $x_1 : A_1, \dots$, and A_n is a type in the context $x_1 : A_1, x_2 : A_2, \dots x_{n-1} : A_{n-1}$:

A_1 type [depending on no assumption]
 A_2 type ($x_1 : A_1$)
...
 A_n type ($x_1 : A_1, x_2 : A_2, \dots x_{n-1} : A_{n-1}$)

Hypothetical judgements introduce functions from A to B :

$$f(x) : B (x : A)$$

It can be read in several ways, for example:

$f(x) : B$ for arbitrary $x : A$
 $f(x) : B$ under the hypothesis $x : A$
 $f(x) : B$ provided $x : A$
 $f(x) : B$ given $x : A$
 $f(x) : B$ if $x : A$
 $f(x) : B$ in the context $x : A$

It is crucial to notice that the notion of function is intensional rather than extensional. Indeed, the meaning of an hypothetical judgement that introduces a function is that whatever element a is substituted for x in $f(x)$, an element $f(a)$ of B results – whereby $f(x)$ is interpreted as a function from A to B . Moreover, the equality of two functions is defined by establishing that substitutions of equal elements of A result in intensionally equal elements of B .⁷

In addition to domains of individuals, an interpreted scientific language requires propositions. Because of the Curry-Howard isomorphism we have

$$\begin{aligned} \text{prop} &: \text{type} \\ \text{prop} &= \text{set} : \text{type} \end{aligned}$$

Propositions are introduced in CTT by laying down what counts as proof of a proposition. Accordingly, a proposition is true if there is such a proof. We write

$$A : \text{prop}$$

to formalize the judgement that A is a proposition. Propositional functions are introduced by hypothetical judgements. Those hypothetical judgements required to introduce propositional functions have the form:

$$B(x) : \text{prop} (x : A)$$

⁷ Cf. Ranta (1994, p. 21), Nordström/Petersson/Smith (1990, chapter 3.3), Primiero (2008, pp. 47-55) and Granström (2011, pp. 77-102).

that reads, $B(x)$ is of the type proposition, provided it is applied to elements of the (type-)set A . The rule by which we produce propositions from propositional functions is the following:

$$\frac{a : A \quad B(x) : \text{prop } (x : A)}{Ba : \text{prop}}$$

And it requires also of the formulation of an appropriate rule that defines the equivalence relation within the type prop :

$$\frac{a=b : A \quad B(x) : \text{prop } (x : A)}{Ba=Bb : \text{prop}}$$

The notion of propositional function as hypothetical judgement allows the (intensional) introduction of subsets by separation:

$$\frac{A : \text{set} \quad B(x) : \text{prop } (x : A)}{\{x : A \mid B(x)\} : \text{set}} \quad \frac{b : A \quad B(b) \text{ true}}{b : \{x : A \mid B(x)\}}$$

This explanation of subsets also justifies the following rules:

$$\frac{b : \{x : A \mid B(x)\}}{b : A} \quad \frac{b : \{x : A \mid B(x)\}}{B(b) \text{ true}}$$

Since this method is based on pre-existent sets that have been constructed by description of their canonical elements, the standard paradoxes of set theory do not appear (such paradoxes also appear in some early formulations of the Lorenzen's method for the construction of sets).⁸

A brief historical interlude: Allow us now a brief historic interlude on the distinction between assertions of the form $a : A$ and of the form $A(a) \text{ true}$ (i.e. $b(a) : A(a)$, where $a : B$, $B : \text{set}$ and $A(x) : \text{prop } (x : B)$). As discussed in Rahman/Clerbout (2015, pp. 145-46) this distinction seems to come close to the reconstruction that Lorenz and Mittelstrass (1967) provide of Plato's notion of *correct naming* in the *Cratylus*. Indeed, in their paper Lorenz and Mittelstrass point out two basic different speech-acts, namely *naming* (ὀνομάζειν) and *stating* (λέγειν). The first speech-act amounts to the act of subsuming one individual under a concept and the latter establishes a proposition about a previously named individual. If the naming has been correctly carried out the (named) individual reveals the concept it instantiates (*names reveal objects for what they are*). *Stating truly* is about the truth of the proposition involving the (correctly) named individual. Moreover, In the context of our own reconstruction *naming* (ὀνομάζειν) corresponds to the assertion that an individual is an element of a given set. That is, it involves judgements of the form

8 Cf. Siegwart (1993).

$$a : A, (A : \text{set})$$

and *stating* (λέγειν) corresponds to building a proposition, such as $A(a)$ from a set, say B , that is

$$A(a) : \text{prop } (a : B).$$

The resulting proposition (in the context of our example) $A(a)$ is true if $A(x)$ can be said of the individual a . In such a case $A(a)$ has been *stated truly*. Furthermore, Lorenz/Mittelstrass (1967, pp. 6-7) argue convincingly that in the context of the *Cratylus* *being correct* and *being true* is to be considered as synonymous. The specialized literature criticized harshly this point of Plato⁹, since, according to the criticism, while truth applies to propositions, it does not apply to individuals. Lorenz/Mittelstrass (1967, pp. 6-12) defended the old master proposing to read these passages as presupposing that in both cases we have a predicate that forms a proposition.¹⁰ Independently of the philological background, if we follow the CTT interpretation, it is not necessary to conclude that in order to claim that both parts of a sentence, namely the noun and the predicative part, are bearers of truth, both must be predicates – provided we move from sentences to assertions. Indeed, let us take once more the letter A , and use it ambiguously again to express both the assertion A *true* (where $A : \text{set}$), and the assertion $A(a)$ *true* (where $A(a) : \text{prop } (a : B)$). From what we presented before on CTT both make perfect sense: whereas A is true if A has an element, and thus asserting that a is an element of A amounts to the truth of A – provided a is indeed such an element; $A(a)$ is true if a is one of those in B of which A can be stated. If we follow this interpretation the point of the correctness is to deploy the *formation rules* presupposed by assertions of the form $a : A$ and of the form $A(a)$. This perspective also suggests to read the expression $A(a)$ *true* (where $A(x) : \text{prop } (x : B)$) as rendering explicit the relation between both, *naming* and *stating*, that Lorenz and Mittelstrass ascribe to Plato – taking exception of their interpretation that both acts involve predicates:

Names, i.e. predicates, are tools with which we distinguish objects from each other. *To name objects or to let an individual fall under some concept is on the other hand the means to state something about objects, i.e. to teach and to learn about objects, as Plato prefers to say.* Lorenz/Mittelstrass (1967, pp. 13).

On our view, these considerations invite to a further generalization: the difference between those speech-acts under discussion in the *Cratylus* is about the difference between categorical

⁹ Viktor Iliewski (2013, pp. 12-13) provides a condensed formulation of this kind of critics:

Socrates next proceeds briefly to discuss true and false speech, with an intention to point out to Hermogenes that there is a possibility of false, incorrect speech. It is a matter of very basic knowledge of logic that truth-value is to be attributed to propositions, or more precisely utterances, specific uses of sentences. Plato's Socrates acknowledges that, but he, somewhat surprisingly, ascribes truth-value to the constituents, or parts of the statements as well, on the assumption that whatever is true of the unit, has to be true of its parts as well. This seems to be an example of flagrant error in reasoning, known as the fallacy of division. Why would Plato's Socrates commit such a fallacy in the course of what seems to be a valid and stable argument? One obvious answer would be that the very theory he is about to expound presupposes the notion of names as independent bearers of meaning and truth, linguistic microcosms encapsulating within themselves both truth-value and reference. In other words, the theory of true and false names has to presuppose that names do not only refer or designate, or even do not only refer and sometimes suggest descriptions, but that they always necessarily represent descriptions of some kind.

¹⁰ Lorenz/Mittelstrass (1967, p. 6):

It follows that a true sentence SP really does consist of the 'true parts' S and P, i.e. $t \in S$ and $t \in P$. In case of a false sentence SP, however, the second part $t \in P$ is false, while the first part $t \in S$ should ex definitione be considered as true, because any sentence is necessarily a sentence about something (Soph. 262e), namely the subject of it. The subject has to be effectively determined, i.e. it must be a thing correctly named, before one is going to state something about it.

assertions, that involve independent types, and hypotheticals that involve dependent ones. This is one lesson that the Plato offers to the reader of his dialogue.

1.4 The fundamentals of the intuitionistic first-order logic in the frame of CTT

In this section we are going to briefly introduce the basic rules of intuitionistic logic in the CTT frame. Our presentation is inspired by the brief description given in Ranta (1991, section 3) of the system of rules introduced by Martin-Löf (1984, pp. 24-25).

1.4.1 Four types of rules

Because of the fact that in this frame propositions are considered as sets, logical operators are defined as set-theoretical ones. Meaning of these operators is established through four types of different rules, namely: formation rules, introduction rules, elimination rules and equality rules :

- *Formation rules.* The explicit inclusion of the formation rules in the inferential system is one of the most distinctive properties of the CTT. Formation rules simultaneously establish the syntax and the basic types, to which logical and non-logical constants of the language in question correspond. More precisely, the formation rules specify under which conditions can we infer that something is a set (a type), and under which conditions can we affirm that two sets are equal. Therefore, since well-formedness includes not only modes of syntactic composition, but, equally, identification of basic corresponding types, we could say that formation rules display at the same time rules of syntactic well-formedness and the meaning specific to a language in question. In fact, any CTT demonstration¹¹ begins by verifying that expression of judgement, which needs to be proved, results from the application of corresponding formation rules. This is a way to implement inside the CTT frame the idea of fully interpreted language: when we read CTT demonstration bottom-up, it displays syntactic and semantic elements of the demonstrated judgement.
- *Introduction Rules.* Introduction rules define the types of the system, establishing the way to form canonical elements and the way to determine if two canonical elements are equal.
- *Elimination rules.* They establish how to define functions (called selectors) on the sets defined by the rules of introduction.
- *Equality rules.* As we mentioned in the previous section, equality rules are linked to the “harmony” between introduction and elimination rules. More precisely, equality rules specify the way in which selectors, defined by elimination rules, work and how to execute their computation in relation to the canonical elements, generated by the introduction rules.

¹¹ We remind to the reader that we use the expression *CTT demonstration* as an abbreviation of the expression *a demonstration in the frame of the constructive type theory*.

1.4.2 The intuitionistic logic in CTT

The Π -operator

(cartesian product of a family of set):

$$\begin{array}{c}
 \begin{array}{c}
 (x : A) \\
 \dots \\
 A : set \quad B(x) : set \\
 \hline
 (\Pi x : A)B(x) : set
 \end{array}
 \quad \Pi F
 \end{array}
 \qquad
 \begin{array}{c}
 (x : A) \\
 \dots \\
 b(x) : B(x) \\
 \hline
 (\lambda x)b(x) : (\Pi x : A)B(x)
 \end{array}
 \quad \Pi I$$

$$\begin{array}{c}
 c : (\Pi x : A)B(x) \quad a : A \\
 \hline
 Ap(c, a) : B(a)
 \end{array}
 \quad \Pi E$$

$$\begin{array}{c}
 (x : A) \\
 \dots \\
 a : A \quad b(x) : B(x) \\
 \hline
 Ap((\lambda x)b(x), a) = b(a) : B(a)
 \end{array}
 \quad \Pi Eq^1$$

$$\begin{array}{c}
 c : (\Pi x : A)B(x) \\
 \hline
 c = (\lambda x)Ap(c, x) : (\Pi x : A)B(x)
 \end{array}
 \quad \Pi Eq^2$$

The introduction rule assumes, as usual, that there is no other occurrence of x in any assumption besides the ones of the form $(x : A)$. The binary function $Ap(x, y)$ (the selector *Application*) is defined by both, the way it is introduced (by the elimination rule) and by the specification of its computation (using equality rules). One can read it as “application of x to y ”. In fact, it is a method to obtain a canonical object $B(a)$, provided $a : A$ (see Martin-Löf (1984, pp.28-29)).

In other words, the computation rule for $Ap((\lambda x)b(x), a)$ establishes that if $b(x)$ is a proof-object for $B(x)$ under the assumption that $x : A$, and given that a is such a proof of A , then its execution yields a proof object $b(a) : B(a)$, such that a replaces x throughout in $b(x)$. For short:

$$Ap((\lambda x)b(x), a) \rightarrow b(a/x)$$

The Π -operator allows to define the universal quantifier and the material implication in the following way:

- $(\forall x : A)B(x) = (\Pi x : A)B(x) : prop$, provided $A : set$ and $B(x) : prop (x : A)$.
- $A \rightarrow B = (\Pi x : A)B : prop$, provided $A : prop$ and $B : prop$.

The Σ operator

(disjoint union of a family of sets)

$$\begin{array}{c}
 (x : A) \\
 \vdots \\
 A : \text{set} \quad B(x) : \text{set} \quad \vdots A \quad b : B(a) \\
 \hline
 (\Sigma x : A) B(x) : \text{set} \quad \Sigma F \quad \hline
 (a, b) : (\Sigma x : A) B(x) \quad \Sigma I
 \end{array}$$

$$\begin{array}{c}
 (x : A, y : B(x)) \\
 \vdots \\
 c : (\Sigma x : A) B(x) \quad (x, y) : C((x, y)) \\
 \hline
 E(c, (x, y) d(x, y)) : C(c) \quad \Sigma E
 \end{array}$$

$$\begin{array}{c}
 (x : A, y : B(x)) \\
 \vdots \\
 a : A \quad b : B(a) \quad d(x, y) : C((x, y)) \\
 \hline
 E(a, b, (x, y) d(x, y)) = d(a, b) : C((a, b)) \quad \Sigma Eq
 \end{array}$$

The expression $E(c, (x, y) d(x, y))$ which occurs as the conclusion of Σ -elimination rule, is informally read as the following computational instruction:

Execute c .

The result of the execution is a canonical element which has the form of the couple (a, b) such that $a : A$ and $b : B$.

Now substitute a and b in the right premise, for x and y respectively.

Thus obtain: $d(a, b) : C((a, b))$.

The execution of $d(a, b)$ will give for result a canonical element e of $C((a, b))$ – it is not difficult to deduce, therefore, that e is a canonical element of $C(c)$ (see Martin-Löf (1984, p. 40)).

The Σ -operator allows to define the existential quantifier and the conjunction in the following way:

- $(\exists x : A) B(x) = (\Sigma x : A) B(x) : \text{prop}$, provided $A : \text{set}$ and $B(x) : \text{prop} (x : A)$.
- $A \wedge B = (\Sigma x : A) B : \text{prop}$, provided $A : \text{prop}$ and $B : \text{prop}$.

In the case of conjunction, we obtain the standard elimination rules from the elimination rules of Σ ,

1. if we decide that C is either A or B , and
2. if we define the projection rules $p(c)$ and $q(c)$, mentioned in the previous section, in the following way: $p(c) \equiv E(c, (x, y) x)$ and $q(c) \equiv E(c, (x, y) y)$.

That is, if we carry out steps 1 and 2, from ΣE we get:

$$c : A \wedge B$$

$$c : A \wedge B$$

$$\frac{}{p(c) : A} \wedge E1$$

$$\frac{}{q(c) : B} \wedge E2$$

Recalling the equality rules, we come to the following computational rules for the execution of $p(c)$ and $q(c)$, where c is constituted by the pair (a, b) such that $a : A, b : B$:

$$p(a, b) \rightarrow a$$

$$q(a, b) \rightarrow b$$

The operator +

(the disjoint union or the co-product of two sets):

$$\frac{A : set \quad B : set}{A + B : set} +F$$

$$\frac{a : A}{i(a) : A + B} +I1$$

$$\frac{b : B}{j(b) : A + B} +I2$$

$$\frac{\begin{array}{ccc} & (x : A) & (y : B) \\ & \dots & \dots \\ c : A + B & d(x) : C(i(x)) & e(y) : C(j(y)) \end{array}}{\mathbf{D}(c, (x)d(x), (y)e(y)) : C(c)} +E$$

$$\frac{\begin{array}{ccc} & (x : A) & (y : B) \\ & \dots & \dots \\ a : A & d(x) : C(i(x)) & e(y) : C(j(y)) \end{array}}{\mathbf{D}(i(a), (x)d(x), (y)e(y)) = d(a) : C(i(a))} +Eq^1$$

$$\frac{\begin{array}{ccc} & (x : A) & (y : B) \\ & \dots & \dots \\ b : B & d(x) : C(i(x)) & e(y) : C(j(y)) \end{array}}{\mathbf{D}(j(b), (x)d(x), (y)e(y)) = e(b) : C(j(b))} +Eq^2$$

The expressions i and j are two new primitive constants (injections), that provide the information that an element of $A+B$ belongs to A or to B , and which of both is actually the case.

The expression $\mathbf{D}(c, (x)d(x), (y)e(y))$ that occurs in the conclusion of the elimination rule is informally read as the following computational instruction:

Execute c ;

if as result of the execution the canonical elements $i(a)$ results, then substitute a for x in $d(x)$;
if as result of the execution the canonical elements $j(b)$ results, then substitute b for y in $e(y)$

If we recall the equality rules, we can also express this in the following way: the execution of $\mathbf{D}(c, (x)d(x), (y)e(y))$, yields a proof-object for $C(c)$ built from proof-object c for $A + B$, and the proof-objects $d(x)$ and $e(y)$ for $C(i(x))$ and $C(j(y))$ respectively, such that $x : A$ and $y : B$. If we display this with two computational rules we obtain:

$$\begin{aligned} \mathbf{D}(i(a), (x)d(x), (y)e(y)) &\rightarrow d(a) \text{ (where } d(a) : C(i(a))) \\ \mathbf{D}(j(b), (x)d(x), (y)e(y)) &\rightarrow e(b) \text{ (where } e(b) : C(j(b))) \end{aligned}$$

The operator $+$ allows to define the disjunction in the following way:

$A \vee B = A + B : prop$, provided $A : prop$ and $B : prop$.

The absurdum $\perp$

$$\begin{array}{ll} \text{----- } \perp F & x : \perp \\ \perp : set & \text{----- } \perp E \text{ (for arbitrary } x) \\ & a : A \end{array}$$

The symbol $\perp$ is nothing else than an empty set, and since its formation rule establishes that it is always a set, the same rule also establishes that $\perp$ is always a proposition. The absurdum does not require any introduction rule. In fact, the elimination rule for $A \rightarrow \perp$ introduces $\perp$. Clearly, it follows from the standard perspective of intuitionistic logic in which $\sim A$ is seen as an abbreviation for $A \rightarrow \perp$. Since there is no specific rule for $\perp$, there is no specific equality rule which would establish the relationship between the introduction and the elimination rules. Finally, the elimination rule expresses the so-called rule of *ex falso sequitur quodlibet*. In fact, a judgement that an x is of type $\perp$ is contradictory because $\perp$ is an empty set. So, from such a judgement we could conclude $a : A$ (see Martin-Löf (1984, pp.65-67)).

Because of the fact that one of the main theses of our work is based on the idea that the copy-cat move in dialogical logic is closely linked to the dynamic version of the notion of definitional equality as deployed by the Π - and Σ -equality rules, we should present systematically CTT-take on definitional and propositional equality.

Due to space limitations, we will not be able to provide a detailed description of the CTT here. In the next section we will limit ourselves by exhibiting the issues, relevant for the objectives of the present work. For more information on this perspective, the reader should examine Martin-Löf (1984), Ranta (1988, 1994), Nordström et al. (1990), Primiero (2008) and Granström (2011).

II.2 Assertions and equality

As mentioned above, whenever a new expression is introduced in CTT, this introduction is deployed by means of what is called a *meaning explanation*. In the case of the introduction of a new type, the meaning explanation consists in (1) describing its canonical elements, (2) providing an algorithm for deciding whether a non-canonical object is of the

given type or not, and (3) give the conditions that allow establishing (or not) the identity of two objects with respect to the given type. The point (3) consists in defining an appropriate equivalence relation. In this way, assertions of the form $a = b : A$, affirm that two objects a and b satisfy the equivalence relation defined for the type A . The assertion $a = b : A$ is also called an *assertion of definitional equality*, since by means of this assertion explicit definitions are **expressed**. This kind of equality is transmitted by reflexivity, symmetry, transitivity, and by substitution of definitionally equals. (Ranta 1994, 52). Even more, real definitions as pointed out by Frege while criticizing Hilbert, might provide the truth of a proposition based on those definitions, but they are *not* bearers of truth. For short, definitional equality is *not* a predicate. The written form $a = b$ is ambiguous unless we make fully explicit the assertion in which it is embedded: at the left of the colon it expresses definitional equality and the right of it expresses a predicate.

When the type is a proposition constituted by logical constants such as conjunction, disjunction and so on, the inferential role of the equalities at the left of the colon is to **harmonize the process of synthesis and analysis that deploys the meaning explanation of** the logical constant at stake. Let us briefly discuss this point in the following section – for a thorough discussion see Rahman/Redmond (2015b).

II.2.1 Proof-objects, equality and their inferential role

In his landmark 1935 article on natural deduction, Gerhard Gentzen points out:

The introductions represent, as it were; the ‘definitions’ of the symbols concerned, and the eliminations are no more, in the final analysis; than the consequence of these definitions. (Gentzen 1969, 80)

The idea behind Gentzen's observation is that an introduction rule in the system of natural deduction exhibits the basis for the assertion of a proposition which contains the logical constant in question, and that the result of using the correspondent elimination rule shows exactly those elements of proposition, that introduction rule characterized as sufficient for its assertion. Dummett calls such form of coordination *harmony*. Briefly, an elimination rule is harmonic if this rule does no more and no less than displaying the consequences of the meaning provided by the introduction rule.¹² If we begin with the introduction rules, one way to see the contribution of the concept of harmony is that while the introduction rules show how to *synthesize*, from certain premises, an expression that contains a given logical constant, the elimination rule shows how to *analyze* the expressions under consideration into exactly those components required by the introduction rule.

The explicit language of the CTT allows **implementing** the condition of harmony by means of definitional equalities. Let us take as an example the case of the conjunction. The proposition $A \wedge B$ (or the set $A \times B$) is explained by establishing that the canonical element of $A \wedge B$ is a couple of proof-objects (a, b) where $a : A$ and $b : B$ – that is to say, where a is a proof-object of A and b is a proof-object of B :

$$\begin{array}{l} a : A \quad b : B \\ \hline (a, b) : A \wedge B \end{array} \quad \wedge\text{-introduction}$$

In order to formulate $\wedge$ -elimination-rule we will use certain operators, called *selectors*, by the means of which new functions could be defined that extract those elements that constitute a

¹² Cf. Stephen Read (2008, 291, note 5).

complex proof-object c (such as $c = (a, b)$). In the case of conjunction, the selectors are the projection functions p and q that yield, respectively, the left and the right side of the couple of proof-objects. Therefore, if c is a proof-object for a conjunction, then $p(c)$ yields the left component of c and $q(c)$ yields its right component.

$$\begin{array}{ccc} c : A \wedge B & & c : A \wedge B \\ \hline p(c) : A & \wedge\text{-}p\text{-elimination} & q(c) : B \end{array} \quad \wedge\text{-}q\text{-elimination}$$

If we know that $c = (a, b)$, then $p(c)$ will restore the right component of c (obtained by the introduction rule). In other words, $p(c) = p((a, b)) = a$, such that $a : A$. Analogously, $q(c)$ restores the right component:

$$\begin{array}{ccc} a : A & b : B & \\ \hline ((a, b)) = a : A & \wedge\text{-left equality} & \end{array} \quad \begin{array}{ccc} a : A & b : B & \\ \hline q((a, b)) = b : B & \wedge\text{-right equality} & \end{array}$$

These are clear examples of how to use the notion of definitional equality, previously mentioned: projection functions p and q are explicitly defined by the inference rule in the way that, given proof-objects a and b , the projection q of (a, b) is definitionally equal to b , according to the equivalence relation which is defined within the type B . The projection p is introduced analogously. These equalities can be seen as displaying the execution of the computation rules mentioned above for the projections p and q in relation to the pair (x, y) , such that $x : A, y : B$

One can also introduce a dual rule, called η , which could be defined in the following way:

$$\begin{array}{c} c : A \wedge B \\ \hline (p(c), q(c)) = c : A \wedge B \end{array} \quad \wedge\text{-}\eta\text{-equality}$$

In the present context these rules could be seen as ensuring that the rules are harmonic. That is to say, if the proof-object c for the conjunction is constituted by the pair (a, b) , then both elimination rules provide an analysis of the proof-object c from which the components a and b result.

II.2.2 Definitional equality and propositional equality

II.2.2.1 The definitional equality

The definitional equality meets the basic conditions of equality:

Reflexivity	Symmetry	Transitivity
$a : A$	$a = b : A$	$a = b : A \quad b = c : A$
$\hline a = a : A$	$\hline b = a : A$	$\hline a = c : A$
$A : \text{set}$	$A = B : \text{set}$	$A = B : \text{set} \quad B = C : \text{set}$

$$\frac{}{A = A : set} \quad \frac{}{B = A : set} \quad \frac{}{A = C : set}$$

Extensionality

$$\frac{A = B : set \quad a : A}{a : B} \quad \frac{A = B : set \quad a = b : A}{a = b : B}$$

In the context of hypothetical judgements the definitional equality between two proof-objects, for example a and c , in the given set, say A , produces the equality between functions $b(a)$ and $b(c)$, which is inferred from the substitution of the variable x in the function $b(x)$ from A to $B(x)$ by **both** a and c .¹³

Substitution

$$\frac{a = c : A \quad b(x) : B(x) (x : A)}{b(a) = b(c) : B(a)} \quad \frac{a = c : A \quad B(x) : set (x : A)}{B(a) = B(c) : set}$$

Hypothetical judgements involving hypothetical equalities yield the following rule:

Definitional equalities between dependent objects

$$\frac{a : A \quad c(x) = b(x) : B(x) (x : A)}{b(a) = c(a) : B(a)}$$

Substitution rules for equalities in the context of hypothetical judgements require equally the general notion of substitution of variables expressed by the following rule:

General substitution rule for hypotheticals

$$\frac{a : A \quad b(x) : B(x) (x : A)}{b(a) : B(a)}^{14}$$

Let us now see how in the frame of CTT the notion of definitional equality and propositional equality relate.

¹³ We present here the rules of substitution for equalities between the elements of a set, but it is equally possible to define analogously equalities between elements of the type *set*, that is equalities between sets. (see Nordström/Petersson/Smith (1990), p. 38).

¹⁴ Actually, this rule does not suffice, because it does not include the case of simultaneous substitutions, that are necessary, for example, in cases as $C(x, y) : set (x : A, y : B(x))$. There are different solutions to this respect, but the most practical is that from a procedural point of view (so dear to dialogical perspective), the one that allows partial substitutions in a context. Returning to our example, if it is the case that $a : A$, we conduct a first application of the substitution to obtain $C(a, y) : set (y : B(a))$, and then a second one to obtain $b : B(a), C(a, b) : set$ (see Nordström/Petersson/Smith (1990), p. 39).

II.2.2.2 Propositional equalities: Predicates of intensional and extensional equality

Now we need to introduce the equality predicate that allows both, expressing equality as a proposition and showing how definitional and propositional equality relate. Actually, it is possible to define two equality predicates, one intensional and one extensional. The proof-object of the latter is independent of the proof-objects of the premises.

- **The predicate of intensional equality**

The predicate of propositional intensional equality, defined over a set, say A , is usually expressed with the notation $\mathbf{Id}(A, a, b)$. Sometimes the notation $a =_A b$ is also used. We will use the first notation, since it shows explicitly that $\mathbf{Id}$ is a relation.

Formation

$$\frac{A : \text{set} \quad a : A \quad b : A}{\mathbf{Id}(A, a, b) : \text{set}}$$

As pointed out by Thompson (1999, p. 110) this formation rule is quite different from the rules for the formation of types we have seen so far. The latter take the form

$$\frac{\dots : \text{type} \quad \dots : \text{type}}{\dots : \text{type}}$$

which means that with these rules alone, we can say what are the types of the system independently of which elements inhabit those types. The rule of $\mathbf{Id}$ -formation above breaks this rule, since it requires that $a : A$ and $b : A$ for $\mathbf{Id}(A, a, b)$ to be a type. This means that the rules generating those types are inextricably bound to inference rules that involve specific proof-objects inhabiting the relevant types.

The basic introduction rule, yields reflexivity and stems from the fact that given the assertion $a : A$, we have then immediately $a = a : A$. From the inferential point of view, this amounts to introducing the reflexive case for $\mathbf{Id}(A, a, a)$.

$$\frac{a : A}{r(a) : \mathbf{Id}(A, a, a)}$$

The proof-object for $\mathbf{Id}(A, a, a)$ is $r(a)$. The inner-structure of this proof-object is nothing more and nothing else than its dependence upon a : what makes the prop $\mathbf{Id}(A, a, a)$, true is an operator, namely r , such that applied to a , yields a .

By using Substitution in sets on $a = b : A$, $\mathbf{Id}(A, a, x) : \text{prop } (x : A)$ and $\mathbf{Id}$ -introduction we obtain $\mathbf{Id}(A, a, a) = \mathbf{Id}(A, a, b) : \text{prop}$.

$$\frac{A : \text{set} \quad a : A \quad x : A^1}{\mathbf{Id-F}}$$

$$\begin{array}{c}
\frac{a = b : A \quad \mathbf{Id}(A, a, x) : \text{prop}}{\mathbf{Id}(A, a, a) = \mathbf{Id}(A, a, b) : \text{prop}} \text{subst}^1 \quad \frac{a : A}{r(a) : \mathbf{Id}(A, a, a)} \mathbf{Id}\text{-introduction} \\
\hline
r(a) : \mathbf{Id}(A, a, b)^{15} \quad \text{set-equality}
\end{array}$$

Hence, these inferential steps yield the non-reflexive case for **Id**:¹⁶:

$$\begin{array}{c}
a : A \quad a = b : A \\
\hline
r(a) : \mathbf{Id}(A, a, b) \quad \mathbf{Id}\text{-introduction}
\end{array}$$

Elimination and Substitution

In relation to the elimination rule, it is related to a generalization of Leibniz's substitution rule, without making use of second-order quantification. The main concept behind Leibniz's substitution rule is that equals with respect to a given predicative identity can be substituted for equals with respect to that identity. Suppose that we have some proof p of a proposition C involving a , and also that we know that $c : \mathbf{Id}(A, a, b)$. We should then be able to infer the proposition C' resulting from replacing (at least) *some* of the occurrences of a in C by b . To capture the idea of substituting *some* of the occurrences we think of C of having the form $C[a/x, a/y]$ in which a replaces two free variables x and y . In fact, more generally, we will also think of the form of C as including cases where also $r(a)$ occurs there. **The point is that every proposition on which a substitution w.r.t. $\mathbf{Id}$ is to be carried out might also include such propositions P where $r(x)$ occurs explicitly – and so those P can be brought to the form $C[a/x, a/y, r(a/x)]$.** This yields the following general rule

$$\begin{array}{c}
a : A \\
b : A \\
c : \mathbf{Id}(A, a, b) \\
C[x, y, z] : \text{set } (x : A, y : A, z : \mathbf{Id}(A, a, b)) \\
d(x) : C[x, x, r(x)] \quad (x : A) \\
\hline
id(c, d) : C[a, b, c] \quad \mathbf{Id}\text{-elimination}
\end{array}$$

What the operator in the expression $id(c, d)$ amounts to is that of applying the identity $c : \mathbf{Id}(A, a, b)$ to d , in our case $d(a)$, and this will return the proof-object $d(a)$. This already suggests the rules that determine when two $id(c, d)$ are equal:

$$\begin{array}{c}
a : A \\
C[x, y, z] : \text{set } (x : A, y : A, z : \mathbf{Id}(A, a, b)) \\
d(x) : C[x, x, r(x)] \quad (x : A) \\
\hline
id(r(a), d) = d(a) : C[a, a, r(a)] \quad \text{equal elements}
\end{array}$$

Let us see now how the elimination rule works with some examples. We start showing that Leibniz's substitution rule can be inferred from the elimination rule. As already mentioned,

¹⁵ Cf. Ranta (1994), p. 53.

¹⁶ Cf. Nordström/Petersson/Smith (1990), p. 57.

the idea of the rule is that given $B(a)$ and $\mathbf{Id}(A, a, b)$, we should be able to conclude $B(b)$ by replacing a with b in $B(a)$.

Leibniz's substitution law

Assume that, $a, b : A$, and $c : \mathbf{Id}(A, a, b)$, and also $p : B(a)$.

The first task to accomplish is to infer from $B(x)$ some expression of the form $C[x, x, r(x)]$ – since this is the form required by the elimination rule. From $B(x)$ we can infer the suitable expression $B(x) \rightarrow B(x)$, with the identity function $\lambda x.x$ as proof-object. Thus, given $\lambda x.x : B(x) \rightarrow B(x)$ (here C (i.e. $B(x) \rightarrow B(x)$) is does not mention $r(a)$) by $\mathbf{Id}$ -elimination we obtain $id(c, (x)\lambda x.x) : B(a) \rightarrow B(b)$, that is : $id(c, (x)\lambda x.x) : B(a) \rightarrow B(b)$.

$$\begin{array}{c}
 \begin{array}{c}
 a : A \quad b : A \quad c : \mathbf{Id}(A, a, b) \\
 \hline
 \end{array}
 \quad
 \begin{array}{c}
 x : B(x) \\
 \hline
 \lambda x.x : B(x) \rightarrow B(x) \quad (x : A) \\
 \hline
 \text{Id-elimination}
 \end{array}
 \quad
 \begin{array}{c}
 \rightarrow I \\
 \hline
 \end{array}
 \end{array}$$

$$\begin{array}{c}
 \hline
 id(c, (x)\lambda x.x) : B(a) \rightarrow B(b)
 \end{array}$$

Making use of the premise $p : B(a)$ and applying implication-elimination we come to $subst(c, p) : B(b)$. Where $subst(c, p)$ is an abbreviation for $ap(id(c, (x)\lambda x.x), p)$.

$$\begin{array}{c}
 \begin{array}{c}
 a : A \quad b : A \quad c : \mathbf{Id}(A, a, b) \\
 \hline
 \end{array}
 \quad
 \begin{array}{c}
 x : B(x) \\
 \hline
 \lambda x.x : B(x) \rightarrow B(x) \quad (x : A) \\
 \hline
 \text{Id-elimination}
 \end{array}
 \quad
 \begin{array}{c}
 \text{III} \\
 \hline
 \end{array}
 \end{array}$$

$$\begin{array}{c}
 \hline
 id(c, (x)\lambda x.x) : B(a) \rightarrow B(b) \quad p : B(a) \\
 \hline
 \text{IIE} \\
 \hline
 ap(id(c, (x)\lambda x.x), p) : B(b)
 \end{array}$$

This leads to the following formulation of Leibniz's rule:

$$\begin{array}{c}
 B(x) : \text{set } (x : A) \quad a : A \quad b : A \quad c : \mathbf{Id}(A, a, b) \quad p : B(a) \\
 \hline
 subst(c, p) : B(b)
 \end{array}$$

$\mathbf{Id}$ is also transmitted by reflexivity, transitivity and symmetry. Since the introduction rule yields reflexivity, we have to show that from the elimination rule we can also obtain symmetry and transitivity

Symmetry: Assume that, $a, b : A$, and $d : \mathbf{Id}(A, a, b)$. Let us consider C to be $e : \mathbf{Id}(A, y, x)$. From $a : A$, we obtain by introduction $r(a) : \mathbf{Id}(A, a, a)$ and this yields by $\mathbf{Id}$ -elimination $id(d, r(a)) : \mathbf{Id}(A, b, a)$. This allows writing down the following rule for symmetry:

$$\begin{array}{c}
 d : \mathbf{Id}(A, a, b) \\
 \hline
 symm(d) : \mathbf{Id}(A, b, a)
 \end{array}$$

where $\text{symm}(d)$ stands for $\text{id}(d, r(a))$.

Transitivity: The main step in order to infer the transitivity of **Id** is to consider C to be $\lambda y.y : \text{Id}(A, y, c) \rightarrow \text{Id}(A, x, c)$. From the premises $d : \text{Id}(A, a, b)$ and $e : \text{Id}(A, b, c)$, the use of **Id**-elimination yields $\text{id}(d, (x)\lambda y.y) : \text{Id}(A, b, c) \rightarrow \text{Id}(A, a, c)$. By Π -Elimination – making use of $e : \text{Id}(A, b, c)$. – we obtain $\text{trans}(d, e) : \text{Id}(A, a, c)$, where $\text{trans}(d, e)$ stands for $\text{ap}(\text{id}(d, (x)\lambda y.y), e)$. This leads us to the following rule:

$$\frac{d : \text{Id}(A, a, b) \quad e : \text{Id}(A, b, c)}{\text{trans}(d, e) : \text{Id}(A, a, c)}$$

- **The predicate of the extensional equality**

The formation rule of the predicate of extensional equality between a and b , defined in the set A ($\text{Eq}(A, a, b)$), is analogous to the predicate of intensional identity:

$$\frac{A : \text{set} \quad a : A \quad b : A}{\text{Eq}(A, a, b) : \text{set}} \quad \text{Eq F}$$

Introduction of Eq

The proof-object eq , that verifies $\text{Eq}(A, a, b)$, does not depend on either a or b . The verification of $\text{Eq}(A, a, b)$ only needs an identity of the arbitrary proof-object. Take eq to be, for example $s = s$.

$$\frac{a = b : A}{\text{eq} : \text{Eq}(A, a, b)} \quad \text{EqI}$$

This rule represents the weakest way in which we could relate ontological equality with a propositional one. In fact, the proof-object eq neither contains, nor depends on any of the proof-objects upon which the assertion of the premise is based. This absence of links between the proof-objects of the premise and the conclusion is even more evident in the following **non-inductively-defined elimination rules**:

Elimination of Eq

$$\frac{c : \text{Eq}(A, a, b)}{a = b : A} \quad \text{EqE}^1$$

In this elimination rule the proof-object c of the premise does not contain anything that could lead to the equality $a = b$ of the conclusion. What is more, since equality expressed by $\text{Eq}(A, a, b)$ is numerical identity, any proof-object c of $\text{Eq}(A, a, b)$ is equal to eq :¹⁷

$$c : \text{Eq}(A, a, b)$$

$$\frac{}{c = eq : \mathbf{Eq}(A, a, b)} \mathbf{EqE}^2$$

It could be shown that with both elimination rules we obtain a substitution rule analogous to the one defined for intensional equality, and, therefore also obtain analogous rules of transmission.¹⁸

$$\frac{\begin{array}{l} a : A \\ b : A \\ c : \mathbf{Eq}(A, a, b) \\ C[x, y, z] : \text{set } (x : A, y : A, z : \mathbf{Eq}(A, a, b)) \\ d(x) : C[x, x, eq] \quad (x : A) \end{array}}{d(a) : C[a, b, c]} \mathbf{EqE}^3$$