

# GEOMETRY OF INTERACTION III : ACCOMMODATING THE ADDITIVES

Jean-Yves Girard  
Laboratoire de Mathématiques Discrètes  
UPR 9016 – **CNRS**  
163, Avenue de Luminy, Case 930  
F-13288 Marseille Cedex 09

*girard@lmd.univ-mrs.fr*

## **Abstract**

The paper expounds geometry of interaction, for the first time in the full case, i.e. for all connectives of linear logic, including additives and constants. The interpretation is done within a  $C^*$ -algebra which is induced by the rule of resolution of logic programming, and therefore the execution formula can be presented as a simple logic programming loop. Part of the data is public (shared channels) but part of it can be viewed as private dialect (defined up to isomorphism) that cannot be shared during interaction, thus illustrating the theme of communication without understanding. One can prove a nilpotency (i.e. termination) theorem for this semantics, and also its soundness w.r.t. a slight modification of familiar sequent calculus in the case of exponential-free conclusions.

## **1 Introduction**

### **1.1 Towards a monist duality**

*Geometry of interaction* is a new form of semantics. In order to understand what is achieved, one has to discuss the more traditional forms of semantics.

#### **1.1.1 Classical model theory**

The oldest view about logic is that of an external observer : there is a pre-existing reality (mathematical, let us say) that we try to understand (e.g. by

proving theorems). This form of *dualism* is backed by the so-called completeness theorem of Gödel (1930), which says that a formula is provable iff it is true in all models (i.e. in all realizations). There is strong heterogeneity in the duality world/observer (or model/proof) proposed by model-theory, since the latter is extremely finite whereas the former is infinite. Hilbert's attempt at reducing the gap between the two actors failed because of the renowned incompleteness theorems, also due to Gödel (1931), whose basic meaning is that infinity cannot be eliminated.

A paradoxical situation arose because Gentzen proved in 1934 a *cut-elimination* theorem for a formulation of logic of his own, *sequent calculus*, yielding the lineaments of an elimination of infinity in the style proposed by Hilbert. The paradox was mainly ideological, since the applications that Gentzen gave of his method to Peano arithmetic (in 1936 and later), make use of very strong forms of infinity... and therefore achieved no elimination of infinity, but for true believers.

The cut-elimination procedure introduced by Gentzen was a pure syntactical rewriting technique for proofs, enabling one to eliminate infinitary notions from finitary theorems, but whose termination could not be proved without even stronger infinitary techniques... Such a hybrid animal had a difficult life, and in particular could not find his status within the narrow duality models/proofs.

### 1.1.2 The semantics of proofs

Of course classical model-theory does not refuse the observer as a minor part of the reality (the same is true for classical physics) : it makes the assumption of an objective reality where notions like *true*, *false*... make sense. A formula  $A$  is true or false in the world, and if I prove  $A$  it is also true that I prove  $A$ . Brouwer, by introducing *intuitionism*, radically changed the classical paradigm, by excluding the external reality and focusing on the interaction of proofs (the first consequence was taxonomical : the creation of *classical* logic, the new name for ordinary simple-minded logic).

Instead of explaining a proof  $\Pi$  of  $A \Rightarrow B$  as *the justification that whenever  $A$  is true then  $B$  is also true*, intuitionism takes the viewpoint of the function which, given as input a proof  $\Sigma$  of  $A$ , yields a proof  $\Pi(\Sigma)$  of  $B$  as output ; the basic example is the identity function, which maps a proof of  $A$  to itself, and which can be seen as a proof of the basic tautology  $A \Rightarrow A$ . This functional viewpoint yields the so-called *semantics of proofs*.

This change of viewpoint should not be too easily styled subjectivistic, even if this was the ideology of Brouwer. This must be seen as a critics of simple-minded realism, analogous to relativity theory, which considers time as the quantified result (and no longer the cause) of interaction. A pure subjectivis-

tic reading is so-called *realizability* which interprets the functions of semantics of proofs as purely syntactical operations, taking the code (e.g. the Gödel-number) of a proof of  $A$  to the code of a proof of  $B$ . But this interpretation is slightly regressive, since the functions involved in the semantics of proofs have a very high degree of *naturality*, which is conspicuous in the following extraordinary fact : whether we take Gentzen's sequent calculus or neighboring paradigms (lambda-calculi, natural deduction), equipped with the rewriting used for eliminating infinity (cut-elimination or  $\beta$ -conversion or normalization), the functions which are induced by proofs are exactly the canonical morphisms of a *cartesian closed category*, a beautiful kind of category (even if one had to wait until 1969 to find the first truly interesting example, *Scott domains*).

It must be observed that intuitionism does not formally contradict the duality models/proofs. For instance there is still a notion of model with a completeness theorem for intuitionistic logic. But the problem is not that one cannot import classical model-theory in the intuitionistic world : it is that it becomes inefficient. Take for instance the notion of *consistency* : in classical logic, in conformity with Hilbert's view, the consistency of  $A$  produces a model of  $A$ , i.e. an object ; but intuitionistic consistency only produces the mockery of an object, a so-called *Kripke model*. The notion of a consistent intuitionistic theory is therefore as ridiculous as the idea of fixing a tire with a horseshoe, nay feeding a horse with leadfree gasoline.

The central point of intuitionism is indeed *constructivity* (which can take the more ideological dress of *constructivism*, that we shall not discuss). Like classical logic (and like any reliable bank), one should be able *on request*, to exhibit something ; no longer a model, but some explicit information<sup>1</sup>. Typically, a proof of a disjunction  $A \vee B$ , could on request be replaced with a proof of  $A$  or with a proof of  $B$ . From this viewpoint, classical logic which allows the principle  $A \vee \neg A$  without being able to tell which one holds, is "inconsistent" (in the common sense, as in the famous saying by king François 1<sup>er</sup> : *women are inconsistent*)<sup>2</sup>.

Technically speaking, this exhibition is made possible by the fact that the cut-free proofs (or the normal terms in  $\lambda$ -calculus) that are the outputs of cut-elimination are proved by very restricted means : typically, in the case of intuitionistic disjunction, the only cut-free way to get  $A \vee B$  is to get it from  $A$  or to get it from  $B$ .

But this explanation would be incomplete, if one were not stressing its "on request" aspect : in real life, a proof of  $A \vee B$  is *never* a proof of  $A$  or a proof

---

<sup>1</sup>In classical logic the model is constructed from the absence of a proof, whereas constructivity tries to extract information from the presence of a proof

<sup>2</sup>Souvent femme varie/Bien fol est qui s'y fie

of  $B$  (only a masochist would state  $A \vee B$  when he knows  $A$ ). In other terms, a proof of  $A \vee B$  is a proof of  $A$  or  $B$ , but we don't know which one ; only on request, we can be more explicit, and tell you which side holds. But this requires a painful work (typically cut-elimination) to replace everywhere the implicit by the explicit.

Here we have to be very precise, *mathematics cannot be explicit* : as soon as one departs from trivial elementary facts like  $2 + 2 = 4$ , mathematics (and all forms of reasoning) become abstract, i.e. *implicit*. This implicit character is conspicuous by the use of *lemmas*, which are combined by transitivity, by means of the rule of *Modus Ponens* : from  $A$  and  $A \Rightarrow B$  deduce  $B$ , rewritten under the form of the *CUT*-rule of Gentzen (we adopt here the formalism of linear logic, see below) :

$$\frac{\vdash \Gamma, A \quad \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta} CUT$$

Without the *CUT*-rule, it is practically impossible to make any deduction (typically, if I know that  $27 \cdot 37 = 999$ , I cannot use the lemma “commutativity of multiplication” to infer  $37 \cdot 27 = 999$ ). Hence the truly paradoxical result of Gentzen is that this essential rule can be eliminated (for pure logic, i.e. predicate calculi), i.e. that mathematics can (within some limits) be explicated... not by a human being, but surely by a machine.

These observations are the starting point for applications to computer science : the choice *L/R* between the two premises of a disjunction can be used to represent a boolean datum, a proof of an implication can represent a computable function, the *CUT*-rule can take care of application of a function to an argument and cut-elimination, suitably implemented, can handle evaluation. The existence of a categorical model basically asserts the robustness of this approach (e.g. independence of any implementation). This is the basis for *functional programming*, more precisely *typed  $\lambda$ -calculus* an experimental language popular among theoreticians.

### 1.1.3 The monist duality

The intuitionistic world replaces the relation *model of  $\neg A$ /proof of  $A$*  of classical logic with the relation *proof of  $A \Rightarrow B$ /proof of  $A$* . This is not a duality : to get a duality, one at least needs an involutive negation, and such a thing does not exist in the intuitionistic world. Linear logic ([G86]) is based on a refined analysis of the categorical semantics of intuitionistic proofs (replacement of Scott domains by *coherent spaces*), and individualizes new logical connectives. The basic point is to remove the rule of contraction of Gentzen (which amounts to saying that from  $A$  one can deduce  $A \wedge A$ ) and also the rule of

weakening (which allows fake hypotheses), in which case the basic symmetries missing in intuitionistic logic are restored ; these symmetries are expressed by the involutive connective of *linear negation*  $A^\perp$ . It is now possible to think of *monist* duality, *proof of  $A^\perp$ /proof of  $A$*  based on the *CUT*-rule, with the essential difficulty that  $A$  and  $A^\perp$  cannot be simultaneously provable (because in the *CUT*-rule,  $\Gamma$  and  $\Delta$  never happen to both empty). In the classical case, the duality is of a very strange kind : a proof of  $A$  and a model of  $\neg A$  cannot simultaneously exist, in particular classical duality cannot account for the difference between two proofs of  $A$ , since there is model of  $\neg A$  on which to compare them. If we want duality to account for the distinction between proofs of the same formula, then the two partners should not be exclusive one of another... The construction of a satisfactory monist duality is therefore delicate (we have to replace *proof of  $A$*  with something slightly more liberal) and has been attacked by various methods, from categorical semantics to game semantics, without yet any definite answer. *Geometry of interaction* is one of the major approaches to this question.

## 1.2 Linear Logic

Coherent semantics ([G86]) is built in analogy to linear algebra. The basic constructions of linear algebra can be mimicked by coherent semantics, and yield the basic linear connectives :

- There is an involutive duality (analogous to vector space duality), which induces *linear negation*  $A^\perp$  ;
- There are operations analogous to the tensor product of vector spaces, and which yield the so-called *multiplicative* connectives  $\otimes, \wp, \multimap$  ;  $\wp$ , which is the disjunction *par*, is the dual of  $\otimes$ , which is the conjunction *times* (the tensors here are strongly non self-dual) ; the *linear implication*  $\multimap$  can be defined by  $A \multimap B = A^\perp \wp B$ . To the multiplicative universe should be attached the dual constants 1 and  $\perp$ .
- There are operations analogous to the direct sum of vector spaces, and which yield the so-called *additive* connectives,  $\&, \oplus$  ;  $\&$ , which is the conjunction *with* is the dual of  $\oplus$ , which is the disjunction *plus* ; here too the sums are strongly non-involutive. To the additive universe should be attached the dual constants  $\top$  and 0.
- There are operations analogous to symmetric tensor algebras, and which yield the so-called *exponential* connectives  $!, ?$ .  $!$  (*of course*) is the dual of  $?$  (*why not*).

The main categorical properties of these connectives are expressed by a certain number of canonical isomorphisms :

- Involutivity of negation (which allow in practice to ignore double negation symbols) ;
- Commutativity of  $\otimes, \wp, \&, \oplus$  ;
- Associativity of  $\otimes, \wp, \&, \oplus$  ;
- Distributivity of  $\otimes$  over  $\oplus$  (and dually of  $\wp$  over  $\&$ ) ;
- Exponentiation isomorphisms  $!(A\&B) \simeq !A\otimes!B$  (and dually  $?(A\oplus B) \simeq ?A\oplus?B$ ) ;
- Neutrality of the constants 1 w.r.t.  $\otimes, \perp$  w.r.t.  $\wp, \top$  w.r.t.  $\&, 0$  w.r.t.  $\oplus$ , together with  $!\top \simeq 1$  and  $?0 \simeq \perp$ .

### 1.3 Linear sequent calculus

Linear logic is organised into a sequent calculus, (in which we can in particular derive the canonical proofs of the isomorphisms just mentioned). Its standard syntax is one-sided, with defined negation (and implication). Sequents are of the form  $\vdash A_1, \dots, A_n$ , where  $A_1, \dots, A_n$  is a sequence of formulas. The rules are organised in three groups :

#### 1.3.1 Identity/negation

$$\frac{}{\vdash A, A^\perp} ID \quad \frac{\vdash \Gamma, A \quad \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta} CUT$$

This group asserts that  $A$  is  $A$ , the only absolute evidence of logic ; this fact is expressed by two rules. In the one-sided calculus, the identity can be seen as the definition of negation.

#### 1.3.2 Structure

$$\frac{\vdash \Gamma, A, B, \Delta}{\vdash \Gamma, B, A, \Delta} X$$

a rule that can be avoided if one considers sequents as multisets instead of sequences. The two other traditional structural rules

$$\frac{\vdash \Gamma}{\vdash A, \Gamma} W \quad \frac{\vdash A, A, \Gamma}{\vdash A, \Gamma} C$$

also called *weakening* and *contraction* are not allowed in linear logic. These rules are essential in classical and (when written in two-sided style) in intuitionistic logics.

### 1.3.3 Logic

$$\frac{\vdash \Gamma, A \quad \vdash \Delta, B}{\vdash \Gamma, \Delta, A \otimes B} \otimes \quad \frac{\vdash A, B, \Gamma}{\vdash A \wp B, \Gamma} \wp \quad \frac{}{\vdash 1} 1 \quad \frac{\vdash \Gamma}{\vdash \perp, \Gamma} \perp$$

The rule for  $\otimes$  concatenates (adds up) the contexts ; if the context is a price to pay (by destruction) to get a formula, a natural price for  $A \otimes B$  is the sum of price for  $A$  and of a price for  $B$ .

$$\frac{\vdash \Gamma, A \quad \vdash \Gamma, B}{\vdash \Gamma, A \& B} \& \quad \frac{\vdash A, \Gamma}{\vdash A \oplus B, \Gamma} \oplus_1 \quad \frac{\vdash B, \Gamma}{\vdash A \oplus B, \Gamma} \oplus_2 \quad \frac{}{\vdash \top, \Gamma} \top$$

The rule for  $\&$  assumes that the two contexts, i.e. prices, are the same ; in other terms a possible price for  $A \& B$  is a price for which one can get any of  $A$  and  $B$ , which means that if both of them are simulataneously available (as expected from a conjunction) for this price, only one of them (up to our choice) will eventually be bought.

$$\frac{\vdash ?\Gamma, A}{\vdash ?\Gamma, !A} ! \quad \frac{\vdash A, \Gamma}{\vdash ?A, \Gamma} d? \quad \frac{\vdash \Gamma}{\vdash ?A, \Gamma} w? \quad \frac{\vdash ?A, ?A, \Gamma}{\vdash ?A, \Gamma} c?$$

These rules are called *promotion*, *dereliction*, *weakening*, *contraction*. The interpretation in terms of prices (or resources) is no longer very convincing for exponentials (especially because of the contraction rule) : this is because the exponential group is a “classical” group, which enables one to instill traditional classical or intuitionistic features inside a calculus which otherwise (in spite of the novelty of its connectives) would not be expressive enough. Indeed exponentials allow weakening and contraction, but no longer as structural (i.e. universal) rules, but as controlled logical rules : indeed the role of exponentials is precisely to control the use of these two rules.

The main connective of linear logic is linear implication (which does not appear in the official right-handed syntax) and which is extremely different from the familiar intuitionistic implication : from  $A$  and  $A \multimap B$  one can still derive  $B \dots$  but we have lost  $A$  in the process. This strong difference with preexistent logical systems is due to the disappearance of the contraction rule which enabled one to produce two copies of  $A$  from one. This fact is responsible for the resource-sensitive character of linear logic. This is also expressed as a

form of *causality* :  $A \multimap B$  is the action “from  $A$  get  $B$ ” which involves, like in physics, a reaction : the destruction of the cause. Now observe that among the rules for exponentials, reappears the contraction rule, but now limited to the holders of the front symbol  $?$ . This is why the interpretation of intuitionistic implication as  $!A \multimap B$  is possible.  $!A$  has the meaning of  $A$  *ad libitum*, i.e. allows us from a single use of  $A$  to its unlimited reuse. One can also say, in analogy with quantum mechanics, that linear logic is about microscopic facts and that the exponentials ensure the transition with the macroscopic world. . . but such tantalizing analogies should be taken with care ; in particular quantum non-determinism has not yet find a definite analogue in terms of linear logic.

The idea of causality is that of a reciprocal annihilation of  $A$  and  $A^\perp$  by the *CUT* -rule ; this goes very well with our idea of a monist (or homogeneous) duality. (Notice the change w.r.t. intuitionistic logic, where there inputs and outputs ; linear logic says that these roles are exchanged by linear negation, and what is called input or output depends on the user). But how does this annihilation work ?

## 1.4 Proof-nets

Although the basic formalisms for explication are basically equivalent, it makes in practice a lot of difference to work with natural deduction (or typed  $\lambda$ -calculus) instead of sequent calculus. If sequent calculus remains the best possible formulation of logic, its cut-elimination procedure spends too much time on bureaucratic details, typically permutation of rules. Natural deduction has the immense advantage (in the absence of disjunction and existence) of being insensitive to the order of rules, and therefore is much more efficient than sequent calculus. The idea was therefore to build a kind of natural deduction for linear logic, the main difficulty being that the tree form of natural deduction (so successful in the absence of disjunction and existence) could no longer be exploited in the presence of an involutive negation, for which the distinction input/output or hypothesis/conclusion no longer makes sense.

*Proof-nets* are this “linear natural deduction”, and their discovery has been a long process : first limited to multiplicatives (without the constants), they have been extended to quantifiers, and more recently to additives [G94]. The main difficulty came because of the graph-like structure of such unfamiliar proofs : given a graph  $\mathcal{G}$  which pretends to be a proof-net (a so-called *proof-structure*), can we decide whether this claim is grounded ? This is the problem of correctness criterions : the solutions found are of the form  $\mathcal{G}$  *successfully passes a certain set of tests*, see [L93] in this volume. Soon after the introduction of multiplicative proof-nets in [G86], the paper [G86A] introduced a major dualist idea, namely that the tests to be passed by a proof of  $A$  are like



*virtual* proofs of its negation  $A^\perp$  ; “virtual” means that these tests (which in general do not represent proofs at all) are “dense” in a set containing all proofs of  $A^\perp$  if such objects were to exist.

Proofs and tests are therefore homogeneous in nature : both can be seen as finite “wirings” of the atoms of formulas, and cut-elimination basically amounted to connect the wires and to “follow the current”. Mathematically such wires are permutations and cut-elimination can be written as a pure operation between permutations, the only role of logic being to guarantee that a certain expression makes sense, i.e. converges. It “only” remained to extend this paradigm to full logic...

## 1.5 Geometry of interaction

The philosophy is the following : elimination of infinity (or better : *explication*) is not the implementation of a clever, but artificial algorithm. On the contrary, there is an intrinsic dynamics of interaction (expressing the physical reciprocal annihilation of two antagonist actors), and logic is here only as a comment on this “physical” phenomenon. The comment is about :

- Specifications of the antagonistic actors in the interaction (in analogy with the shape of plugs in electricity which are not responsible for the passage of the current, but which limits in principle the plugging of an *acceptor of 220V* to a *giver of 220V*) ; the specification is therefore a formula.
- The building of a step-by-step justification of the specification ; this justification is what we usually see as a syntactical proof.

The problem was to find the right kind of mathematical objects. The multiplicative case was handled by means of finite permutations. A permutation of  $n$  can be represented by an isometry of the finite-dimensional Hilbert space  $\mathcal{C}^n$ <sup>3</sup>. The need to represent the contraction rule for the connective “?” (and whose dynamics is duplication) leads one to replace finite dimension by “continuous” dimension. In other terms this means that the general case will be handled by means of certain operators (technically : partial isometries) of the standard separable Hilbert space  $\mathcal{H}$ . Indeed we are interpreting proofs by objects of a  $C^*$ -algebra (acting on the Hilbert space, but consisting only of very specific operators). This  $C^*$ -algebraic aspect is not technically essential, but it was the constant leading intuition for the whole program, and without a backbone made of solid mathematics (presumably not yet used in a significant way) nothing would have existed.

---

<sup>3</sup>The appendix contains the basic definitions relevant to this subsection

Cut-free proofs are interpreted by such operators, the meaning being that of a static linear input/output machinery. Typically the linear logic axiom  $\vdash A, A^\perp$  is represented by an extension cord, whose abstract definition is to have two extremities complementarily labelled or shaped (here  $A, A^\perp$ ) and whose physical action is to transfer the input in  $A$  as an output to  $A^\perp$  and vice versa : such an operator can be written as a  $2 \times 2$  anti-diagonal matrix. In general all logical rules are interpreted by isomorphisms of  $C^*$ -algebras, and in case of binary rules, summations. Typically, in [G88] the interpretation was based on the  $*$ -isomorphisms induced by an isometry of  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$  and  $\mathcal{H} \otimes \mathcal{H}$  onto  $\mathcal{H}$ , see subsections 5.8 and 5.9 in appendix, see also [DR93], this volume.

To achieve a dynamical effect, we must give inputs to our operators ; this is achieved in presence of *CUT* : the general paradigm of interpretation is that of a pair  $(U, \sigma)$  (with  $\sigma = 0$  in the cut-free case). The partial isometry  $\sigma$  is hermitian (i.e. it is a partial symmetry) and expresses a feedback. This is our way of saying that *CUT* is a physical plugging. Take the example of a cut between  $\vdash \Gamma, A$  and  $\vdash A^\perp, \Delta$  : the two operators  $V$  and  $W$  enjoying these specifications are indeed square matrices (respectively labelled by the indices  $\Gamma, A$  and  $A^\perp, \Delta$ ). We can join the sets of indices into  $\Gamma, A, A^\perp, \Delta$ , and form  $U$  (which corresponds to a “disjoint sum” of  $V$  and  $W$ ). Now we can introduce the partial symmetry  $\sigma$  which exchanges the two opposite labels and is zero everywhere else. From the pair  $(U, \sigma)$  it is possible to express the feedback, namely that the inputs coming through  $A$  and  $A^\perp$  are equal to the outputs coming out of the same channels, but flipped by  $\sigma$ . This amounts to writing a linear equation (whose parameters are the inputs labelled by the remaining free plugs  $\Gamma, \Delta$ ). For instance assume that  $U$  maps a direct sum  $\mathcal{H} \oplus \mathcal{H}$  to itself, and that  $\sigma$  is a feedback on the first coordinate (i.e.  $\sigma^2$  is the projection on the first coordinate) ; then the problem is, given  $x \in \mathcal{H}$ , to find  $y, a \in \mathcal{H}$  such that :

$$U(\sigma(a) \oplus x) = a \oplus y$$

A sufficient condition for a solution is the invertibility of  $1 - \sigma U$ , in which case the *execution formula*

$$RES(U, \sigma) = (1 - \sigma^2)U(1 - \sigma U)^{-1}(1 - \sigma^2)$$

expresses the solution, i.e. the input/output dependency of the remaining plugs (i.e.  $y$  as a function of  $x$ ), see subsection 5.7 in appendix.

A stronger existence condition is the *nilpotency* of  $\sigma U$ , i.e. that some power  $(\sigma U)^n$  is zero, a condition which is experimentally true for all pairs  $(U, \sigma)$  arising from logical proofs<sup>4</sup>. In this case the central part of the execution

---

<sup>4</sup>For non-logical systems like *pure  $\lambda$ -calculus* there is still a notion of *weak nilpotency*, see [G88A, MR91]

formula can be written as a finite sum :

$$U(1 - \sigma U)^{-1} = U + U\sigma U + U\sigma U\sigma U + \dots$$

the length of the sum being equal to the *order of nilpotency*  $n$  of  $\sigma U$ . Observe that  $n$  can be extremely big : the basic way to get a high order of nilpotency  $n$  is to make  $n$  cuts with identity axioms, which is not very surprising ! But the logical rules, interpreted by  $*$ -isomorphisms will alter the pattern, since if  $*$ -isomorphisms cannot affect invariants like orders of nilpotency etc., they surely can dramatically affect the apparent size (i.e. the description) of objects : think for instance that the sum  $U_1 + \dots + U_{100}$  of 100 isometric copies of  $U$  can easily be recovered from a single operator of the form  $U \otimes 1$  which has a more compact definition (this is the basic idea behind the interpretation of !). In fact this change of size is so dramatic that the order of nilpotency can hardly be predicted from the pair  $(U, \sigma)$ , and that the nilpotency of all pairs  $\Pi \bullet \sigma$  coming from proofs in standard logical systems cannot be proved within usual mathematics : this nilpotency is related to the termination of cut-elimination, and therefore implies the consistency of various systems (in consequence, by incompleteness it cannot be proved without a heavy use of infinitary notions). The *execution formula* does not quite correspond to syntactical cut-elimination, but it is not too far ; in particular for proofs of sufficiently simple formulas, the correspondence is exact and that's enough. The interest of this approach for computer science was later confirmed by its applications to optimal reduction in  $\lambda$ -calculus due to Gonthier ([GAL92]).

## 1.6 The case of additives

The original paper [G88] only dealt with multiplicatives, exponentials and quantifiers of any order, which was enough to modelize extant typed  $\lambda$ -calculi. However some essential elements were missing, typically the treatment of additive connectives. The situation remained the same for several years, until a satisfactory extension of *proof-nets* to the additive case was found [G94]. The main novelty consists in assigning boolean *weights* to the links of a proof-nets. Since geometry of interaction uses a  $C^*$ -algebraic formulation and boolean algebras are basically commutative  $C^*$ -algebras, there is no major obstacle to this extension.

The main difficulty arises in the interpretation of the  $\&$ -rule, typically in presence of a context. For instance, if  $U$  and  $V$  are operators on  $\mathcal{H} \oplus \mathcal{H}$  corresponding to proofs of  $\vdash C, A$  and  $\vdash C, B$ , we must "merge"  $U$  and  $V$  in order to represent the result of the application of the  $\&$ -rule. If we see  $U$  and  $V$  as  $2 \times 2$ -matrices  $(U_{ij})$  and  $(V_{ij})$ , this merge is performed in quite different ways, depending on the index, typically :

- $U_{22}$  and  $V_{22}$  (which correspond to the formulas  $A$  and  $B$ ) can be merged into a single operator  $W_{22}$  (corresponding to the formula  $A \& B$ ), by means of an isometry  $\varphi$  of  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$ , as in the multiplicative case. The isometry  $\varphi$  is used to relate  $A \& B$  to its components  $A$  and  $B$  : it is public knowledge that the coefficient  $W_{22}$  of *any* proof of  $\vdash C, A \& B$  comes from coefficients  $U_{22}$  and  $V_{22}$  by means of  $\varphi$ .
- But the case of  $U_{11}$  and  $V_{11}$  is more delicate : these coefficients correspond to the common context  $C$ . A plain summation (i.e.  $W_{11} = U_{11} + V_{11}$ ) would definitely be too brutal, erasing the fact that  $U_{11}$  is related to  $A$  and not to  $B$ . Therefore  $U_{11}$  and  $V_{11}$  must also be merged by means of an isometry  $\varphi'$  of  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$ . The problem arises from the fact that the formula  $C$  is used to label  $U_{11}$  and  $V_{11}$  and  $W_{11}$ , and so that nothing in  $C$  indicates that we have to look for a decomposition along  $\varphi'$  ; worse, the main connective of  $C$  (e.g. if  $C$  begins with a “&”) suggests another decomposition, along say  $\varphi$ . These two decompositions should be simultaneously possible.
- The trick is to introduce an auxiliary space : in [G88] a cut-free proof  $\Pi$  of a  $n$ -ary sequent was interpreted by a  $n \times n$  matrix with entries in a  $C^*$ -algebra  $\Lambda^*$ . This operator could be seen as acting on a  $n$ -ary direct sum  $\mathcal{H} \oplus \dots \oplus \mathcal{H}$ . Now, our operators will still act on a  $n$ -ary direct sum, but of  $n$  spaces  $\mathcal{H} \otimes \mathcal{H}'$  where  $\mathcal{H}'$  is a Hilbert space, seen as a space of *auxiliary messages*<sup>5</sup>. The conflict between the  $\varphi$  and the  $\varphi'$  decomposition which occurs in the case of  $W_{11}$  is solved by taking for  $\varphi$  and  $\varphi'$  the isometries of  $(\mathcal{H} \otimes \mathcal{H}') \oplus (\mathcal{H} \otimes \mathcal{H}')$  into  $\mathcal{H} \otimes \mathcal{H}'$  respectively induced by an isometry of  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$  and  $\mathcal{H}' \oplus \mathcal{H}'$  into  $\mathcal{H}'$ . The space  $\mathcal{H}'$  is therefore specifically used to handle additive merges.
- But this is not enough, since the same  $C$  can serve several times as a context to an additive rule : think of a sequent  $\vdash C, A \& B, A' \& B'$  which involves the merge of four coefficients : the two possible orders of performance of the  $\&$ -rules induce two alternative merges of the four coefficients. No commutation trick (like the distinction between  $\mathcal{H}$  and  $\mathcal{H}'$ ) can help us any longer. But it is easy to see that the two solutions proposed are isomorphic, i.e. that they are equal up to a change of the auxiliary messages, i.e. up to an isometry of  $\mathcal{H}'$ . This introduces the most important idea, the real novelty of this paper : everything dealing with  $\mathcal{H}'$  is up to isomorphism... This means that everything related to  $\mathcal{H}'$  is *private*, and we therefore speak of a *dialect*. The dialect

---

<sup>5</sup>Of course  $\mathcal{H}'$  is isomorphic to  $\mathcal{H}$ , but we prefer to keep different names here, since both spaces are used in a strongly different spirit

(which behaves to some extent like bound variables in traditional syntax) is up to isomorphism, what we express by introducing an equivalence relation : being *variants*. All constructions are compatible with that equivalence<sup>6</sup>. When we interpret a conjunctive rule,  $\otimes$  or  $\&$ , a common dialect has to be produced, but we should beware of accidental matchings of the respective dialects, exactly like in usual syntax, we have to prevent accidental collision of bound variables. The only possibilities are to make generic operations, tensorization of dialects or summation. On the other hand  $\mathcal{H}$  represents (shared) *channels* of communication. The paradigm of interaction is therefore *communication without understanding through shared channels*.

- As in [G88] the *execution* formula does not quite correspond to syntactical normalisation ; in other terms certain operations (typically erasings) are not performed. But the soundness of the execution formula w.r.t. cut-elimination still holds under a reasonable restriction on the proven formula. Notice that the syntax has to be adapted to prove this result (introduction of the *b-calculus*).

## 1.7 Closing the system

The interpretation given here, although perhaps not definitive, omits no logical connective. But there is still an essential lingering problem : the problem of “closing the system”. This means being able to build a duality between proofs of  $A$  and proofs of  $A^\perp$ . We already noticed that the notion of proof has to be liberalized to something wider (like the tests for  $A^\perp$  in the case of multiplicative proof-nets) but homogeneous to proofs. This is for instance what we do in our definition of *weak types*, see definition 12. But the problem is that, among the elements of the weak type  $A^\bullet$  interpreting  $A$  there is no way to distinguish those objects which are interpretations of proofs. Surely the notion of *orthogonality* introduced in definition 11 and central in the notion of a weak type, is too lax. What is missing is that  $U \perp V$  is completely symmetrical in the partners, and that there is no way to tell the wheats (real proofs) from the tares (tests). This could be fixed by defining (but how ?) an output  $\langle U, V \rangle$  of the interaction, sufficiently antisymmetric so that a distinguished value -say 1- for  $\langle U, V \rangle$  would exclude the same value for  $\langle V, U \rangle$ . It would of course remain to prove completeness, namely that if  $\langle U, V \rangle = 1$  for all  $V \in A^{\perp\bullet}$ , then  $U$  is the interpretation of a proof of  $A$ , a widely open program...

---

<sup>6</sup>But one, namely the promotion rule for ! which strongly resists to the proof-net spirit : in [G86], the promotion rule is accommodated with a “box”, i.e. treated as a global entity ; this globality is perhaps inherent to this rule and technically expressed by a failure of the variance principle

This is where geometry of interaction merges with *game semantics* : the weak orthogonality between  $U, V$  means that  $U, V$  can be seen as strategies for the two players inside the same game, and our output  $\langle U, V \rangle$  is the result of the game, won by  $U$  when the result is 1. Therefore the problem of closing the system is closely related to the search for a complete game semantics for linear logic... (the connection between linear logic and game semantics was initiated by Blass, see [B92]).

## 1.8 Notes on the style

Geometry of interaction is most naturally handled by means of  $C^*$ -algebras ; this yields surely more elegant proofs, but it obscures the concrete interpretation. So we prefer to follow a down to earth description of the interpretation. An unexpected feature will help us : the  $C^*$ -algebras used can in fact be interpreted in terms of *logic programming*, since the basic operators are very elementary PROLOG programs, and composition is resolution ! This ultimate explanation of the dynamics of logic in terms of (some theoretical) logic programming cannot be without consequence : we could for instance try to use linear logic as a typing (i.e. specification) discipline for such logic programs... this promising connection is therefore a strong reason to remain concrete. In appendix we give some hints as to the  $C^*$ -algebraic presentation, enough to understand the various allusions to the Hilbert space that are scattered in the main text.

Although our main inspiration was our (yet unwritten) work on additive proof-nets, proof-nets will not at all appear below : first this would make too many simultaneous novelties, second certain details ( $b$ -proof-nets, simultaneous treatment of additives and quantifiers) have not yet been fixed.

Finally, we shall give many examples, especially in the section devoted to exponentials ; it must not be inferred from the elementary character of these examples that the global construction is trivial : like in  $\lambda$ -calculus, all atomic steps are elementary, but the combination of few such steps easily becomes... explosive. Concerning the style, we are deeply indebted to the referee who read the first version of this work in details and suggested many modifications that should increase the legibility of the paper.

## 2 The algebra of resolution

### 2.1 Resolution

By a *term language*  $T$ , we mean the set of all terms  $t$  that can be obtained from variables  $x_1, \dots, x_n$  by means of a finite stock of function letters ; we assume

that  $T$  contains at least one constant (0-ary function letter), so that the set of ground (i.e. closed) terms is non-empty. By a *language*  $L$ , we mean the set of all atoms  $pt_1 \dots t_n$ , where  $p$  is a  $n$ -ary predicate letter varying through a finite set of such symbols, given with their arity. An important case is when the  $n$  predicates of  $L$  are of the same arity  $m$ , in which case we shall use the notation  $T^m \cdot n$  for  $L$ .

### Remark 1

The basic need of geometry of interaction is to have two constants  $g$  and  $d$  and a binary function  $\odot$ . The predicate letters correspond to the indices of matrices, i.e. a  $n \times n$  matrix makes use of  $p_1, \dots, p_n$  (in the sequel we shall rather use formulas as indices :  $p_A, p_B, \dots$ , which supposes that they are pairwise distinct, a standard bureaucratic problem, solved by replacing formulas by occurrences etc.). These predicates will be binary, although in the presence of a binary function, the arity of the predicates hardly matters.

### Definition 1 (Unification)

Let  $L$  be a language ; we say that two expressions (terms or atoms)  $e$  and  $e'$  are unifiable when there is a substitution (called a unifier of  $e$  and  $e'$ )  $\theta$  of terms for the variables occurring in  $e$  and  $e'$  such that  $e\theta = e'\theta$ .

### Remark 2

In such a case, there is a *most general unifier* (mgu), i.e. a unifier  $\theta$  of  $e$  and  $e'$  such that any other unifier  $\theta'$  can be written as the composition  $\theta\theta''$  of  $\theta$  with a substitution  $\theta''$ . Most general unifiers are unique up to renaming of variables.

### Definition 2 (Clauses)

Let  $L$  be a language ; by a (rudimentary) clause in  $L$ , we mean a sequent  $P \mapsto Q$ , where  $P$  and  $Q$  are atoms of  $L$  with the same variables.

### Remark 3

Clauses are considered up to the renaming of their variables : we do not distinguish between  $p(x) \mapsto q(x)$  and  $p(y) \mapsto q(y)$  (but we distinguish between  $p(x \odot y) \mapsto q(x \odot y)$  and  $p(y \odot x) \mapsto q(x \odot y)$  ; in other words the variables of clauses are bound, and other notations like  $\forall x_1 \dots x_n P \Rightarrow Q$  could be considered.

### Definition 3 (Resolution)

If  $P \mapsto Q$  and  $P' \mapsto Q'$  are clauses, then we can assume w.l.o.g. that they have no variable in common and

- either  $Q$  and  $P'$  are not unifiable, in which case we say that resolution of the clauses  $P \mapsto Q$  and  $P' \mapsto Q'$  (in this order) fails

- or  $Q$  and  $P'$  have a mgu  $\theta$ , in which case resolution succeeds, with the resolvent  
 $P\theta \mapsto Q'\theta$  ; in that case we introduce the notation

$$(P \mapsto Q) \cdot (P' \mapsto Q') = P\theta \mapsto Q'\theta.$$

#### Remark 4

Let us introduce the formal clause 0, and extend resolution by setting  $(P \mapsto Q) \cdot (P' \mapsto Q') = 0$  when unification fails ; then it turns out that resolution is associative. All the clauses  $P \mapsto P$  are idempotent. Finally the operation defined on clauses by  $(P \mapsto Q)^* = Q \mapsto P$  is an anti-involution :  $(R \cdot R')^* = R'^* \cdot R^*$ .

#### Remark 5

In case  $L$  has a single (let us say : binary) predicate symbol  $p$ , the clause  $I = p(x, y) \mapsto p(x, y)$  is neutral w.r.t. resolution : if  $x, y$  are not free in  $r, s, t, u$  then  $p(x, y)$  unifies with  $p(r, s)$  by means of the mgu  $\theta(x) = r, \theta(y) = s$ , hence  $I(p(r, s) \mapsto p(t, u)) = p(x, y)\theta \mapsto p(t, u)\theta = p(r, s) \mapsto p(t, u) \dots$

## 2.2 The algebra $\lambda^*(L)$

### Definition 4 (The algebra)

Let  $\lambda^*(L)$  be the set of all (finite) formal linear combinations

$$\sum \alpha_i \cdot (P_i \mapsto Q_i),$$

with the scalar  $\alpha_i$  in  $\mathcal{C}$  ;  $\lambda^*(L)$  is obviously equipped with

- A structure of complex vector space.
  - A structure of complex algebra, the multiplication being extended by bilinearity from resolution :
- $$(\sum \alpha_i \cdot (P_i \mapsto Q_i))(\sum \beta_j \cdot (R_j \mapsto S_j)) = \sum \alpha_i \beta_j \cdot (P_i \mapsto Q_i)(R_j \mapsto S_j)$$
- An identity : for instance the identity of  $T^2 \cdot n$  is  $\sum p_i(x, y) \mapsto p_i(x, y)$  ( $x, y$ , distinct variables) : easy consequence of remark 5.
  - An (anti-)involution defined by  $(\sum \alpha_i \cdot (P_i \mapsto Q_i))^* = \sum \overline{\alpha_i} \cdot (Q_i \mapsto P_i)$ .

In other terms  $\lambda^*(L)$  bears all the features of a  $C^*$ -algebra, (see definition 26 in appendix) but for the norm features. Although the uses of this fact are quite limited<sup>7</sup>, it is of interest to make a  $C^*$ -algebra out of  $\lambda^*(L)$ . This is done in appendix, see subsection 5.6.

---

<sup>7</sup>The only known utilization of the Hilbert space in geometry of interaction is due to Danos & Regnier [DR93].



## 2.3 The execution formula

### Definition 5 (Loops)

Let us fix the language  $L$ . We adapt the main definitions and notions coming from [G88].

- A message is any finite sum  $\sum(P_i \mapsto P_i)$ , with the  $P_i$  and  $P_j$  not unifiable for  $i \neq j$ . A message is therefore a projection (see subsection 5.5) of  $\Lambda^*(L)$ , and messages commute with each other. We abbreviate the message  $P_i \mapsto P_i$  in  $P_i$ .
- A wiring is any finite sum  $\sum(P_i \mapsto Q_i)$ , with the  $P_i$  and  $P_j$  not unifiable for  $i \neq j$  and the  $Q_i$  and  $Q_j$  not unifiable for  $i \neq j$ ; each of the clauses  $(P_i \mapsto Q_i)$  is called a wire. A finite wiring is therefore a partial isometry (see subsection 5.5) of  $\Lambda^*(L)$ . (In [G88], wirings were called “observables”.) If  $w$  and  $w'$  are wirings then  $ww'$  is a wiring.
- If  $m$  is a message and  $w$  is a wiring, there is a (non-unique) message  $m'$ , such that  $m'w = wm$  : take  $m' = wmw^*$  (PROOF:  $wmw^*w = ww^*wm = wm$ .  $\square$ ) ; in other terms wirings propagate messages.
- A loop  $(U, \sigma)$  is a pair of wirings such that  $\sigma$  is hermitian, i.e.  $\sigma = \sigma^*$  (in particular  $\sigma^2$  is a projection and  $\sigma^3 = \sigma$ ).
- A loop converges when  $\sigma U$  is nilpotent, i.e. when  $(\sigma U)^n = 0$  for some  $n$ . The execution of the loop is the element

$$EX(U, \sigma) = U(1 - \sigma U)^{-1} = U \sum_{k=0}^{n-1} (\sigma U)^k$$

and the result of the execution is defined as

$$RES(U, \sigma) = (1 - \sigma^2)U(1 - \sigma U)^{-1}(1 - \sigma^2).$$

The output is a wiring, whereas the execution is not a wiring. Observe that  $U = V(1 + \sigma V)^{-1}$ , with  $V = EX(U, \sigma)$ .

### Remark 6

Wirings correspond to very specific PROLOG programs. (Here we do not refer to the actual language PROLOG, but rather to the idea of resolution, independently of any implementation). Each wire  $P_i \mapsto Q_i$  in  $\sum(P_i \mapsto Q_i)$  can be seen as a clause in a program : the program consists in the clauses  $P_i \mapsto Q_i$ . Such programs are very peculiar :

- The *tail* (i.e. the body) of the clause consists of a single literal,

- The same variables occur in the head and in the tail,
- The heads are pairwise non-unifiable,
- The tails are pairwise non-unifiable.

But one has to be (slightly) subtler to relate execution in our sense to the execution of a logic program. In order to interpret a loop, we systematically duplicate all predicate symbols  $p, q, r, \dots$  into pairs  $(p^-, p^+)$  ( $p^-$  for inputs,  $p^+$  for outputs) ; therefore an atom  $P$  receives two interpretations  $P^-$  and  $P^+$ . The pair  $(U, \sigma)$  yields the following program, consisting in the clauses

- Clauses  $P^- \mapsto Q^+$  for each wire  $P \mapsto Q$  in  $U$ ,
- Clauses  $P^+ \mapsto Q^-$  for each wire  $P \mapsto Q$  in  $\sigma$ .

The execution of  $(U, \sigma)$  consists in finding all clauses  $P^\epsilon \mapsto Q^{\epsilon'}$  which are consequences of this program by means of resolution. The result of the execution consists in retaining only those clauses  $P^- \mapsto Q^+$  that cannot be unified (by prefixing and/or suffixing) with a clause coming from  $\sigma$ .

### Exercise 1

*By suitable modifications of the language  $L$ , express the execution formula as a logic programming problem such that*

- $L$  has only three unary predicates  $e, c, s$
- $(U, \sigma)$  is interpreted by only four kinds of clauses
  - clauses  $e(t) \mapsto c(u)$ ,
  - clauses  $e(t) \mapsto s(u)$ ,
  - clauses  $c(t) \mapsto c(u)$ ,
  - clauses  $c(t) \mapsto s(u)$ ,

*in such a way that the result of the execution corresponds to the consequences of the program (by means of resolution) of the form  $e(t) \mapsto c(u)$ .*

## 3 The interpretation of MALL

### 3.1 Laminated wirings

Strictly speaking, the interpretation can be made within a fixed algebra  $\lambda^*(L)$ , for instance with one constant, one binary function and one unary predicate... but this not very user-friendly. In practice we shall need two constants  $g$  and

$d$  and one function letter  $\odot$ , together with a variable number of predicates to reflect the structure of *sequents*, which have a variable number of formulas. Also these predicates will be chosen *binary*, for reasons that will be explained below. In the sequel  $T$  will be the term language  $\{g, d, \odot\}$ , and  $T^2 \cdot n$  will be the language built from the terms of  $T$  by means of  $n$  binary predicates,  $p_1, \dots, p_n$ . We shall use the following notational convention for terms :  $t_1 \dots t_k$  is short for  $t_1 \odot (\dots \odot (t_{k-1} \odot t_k) \dots)$ , so that  $tuv$  is the same as  $t(uv)$ .

**Definition 6 (Variants)**

Let  $U$  and  $V$  be wirings in  $\lambda^*(T^2 \cdot n)$  ;  $U$  and  $V$  are said to be variants if there exist wirings  $W$  and  $W'$  of  $\lambda^*(T^2 \cdot n)$  such that :

- $U = W^* V W$  and  $V = W'^* U W'$
- $W$  (resp.  $W'$ ) is a sum of wirings of the form  $\sum_{i=1}^n p_i(x, t) \mapsto p_i(x, u)$  with  $x$  not free in  $t, u$  ; this means that  $W$  and  $W'$ , as operators, are of the form  $Id \otimes Z$ .

The notion of being variants is clearly an equivalence relation, noted  $\sim$ .

**Proposition 1**

If  $U \sim V$  and  $\sigma$  is a sum of wires of the form  $p(t, x) \mapsto q(u, x)$  ( $x$  not occurring in  $t, u$ ), then :

- $U\sigma$  is nilpotent iff  $V\sigma$  is nilpotent ;
- In this case,  $EX(U, \sigma) \sim EX(V, \sigma)$  and  $RES(U, \sigma) \sim RES(V, \sigma)$ .

PROOF: immediate.  $\square$

**Definition 7 (Changing dialects, see remark 7 below)**

Let  $U$  be a wiring in  $\lambda^*(T^2 \cdot n)$  ;

- $\otimes_1(U)$  is defined as follows : every wire  $p(t, u) \mapsto q(t', u)$  of  $U$  is replaced with the wire  $p(t, uy) \mapsto q(t', uy)$ , where  $y$  is a fresh variable.
- $\otimes_2(U)$  is defined as follows : every wire  $p(t, u) \mapsto q(t', u)$  of  $U$  is replaced with the wire  $p(t, xu) \mapsto q(t', xu)$ , where  $x$  is a fresh variable.
- $\&_1(U)$  is defined as follows : every wire  $p(t, u) \mapsto q(t', u)$  of  $U$  is replaced with the wire  $p(t, ug) \mapsto q(t', ug)$  .
- $\&_2(U)$  is defined as follows : every wire  $p(t, u) \mapsto q(t', u)$  of  $U$  is replaced with the wire  $p(t, ud) \mapsto q(t', ud)$  .

**Proposition 2**

- $\otimes_1(U) \sim \otimes_2(U)$
- $\&_1(U) \sim \&_2(U)$
- $\otimes_1(U) \sim \otimes_1(\otimes_1(U))$

PROOF: let us verify the third fact : let  $W = \sum_{i=1}^n p_i(x, (yz)z) \mapsto p_i(x, yz)$  and

$W' = \sum_{i=1}^n p_i(x, yz_1z_2) \mapsto p_i(x, (yz_1)z_2)$  ; then  $\otimes_1(U) = W^* \otimes_1(\otimes_1(U))W$  and  $\otimes_1(\otimes_1(U)) = W'^* \otimes_1(U)W'$ .  $\square$

**Definition 8 (Laminated wirings)**

A wiring  $U$  in  $\lambda^*(T^2 \cdot n)$  is said to be laminated if it is a sum of wires the form  $p(t, u) \mapsto q(t', u)$  and if  $U \sim \otimes_1(U)$ .

**Remark 7**

A proof  $\Pi$  of a  $n$ -ary sequent using  $m$  cuts will be interpreted by a pair  $(\Pi^\bullet, \sigma)$  of laminated wirings in  $\lambda^*(T^2 \cdot (2m + n))$ . We can see  $\Pi^\bullet$  as an operator on the Hilbert space

$(\mathcal{H} \otimes \mathcal{H}) \oplus \dots \oplus (\mathcal{H} \otimes \mathcal{H})$ , a  $n$ -ary direct sum of spaces  $\mathcal{H} \otimes \mathcal{H}$ , which can also be seen as

$(\mathcal{H} \oplus \dots \oplus \mathcal{H}) \otimes \mathcal{H}$ . Now, in  $(\mathcal{H} \oplus \dots \oplus \mathcal{H}) \otimes \mathcal{H}$  the two components are treated in a very different spirit :

- The first component  $(\mathcal{H} \oplus \dots \oplus \mathcal{H})$  is seen as a space of *shared* messages, the *channels* ; typically the decomposition of  $(\mathcal{H} \oplus \dots \oplus \mathcal{H})$  into its summands is part of this public knowledge.
- The second component  $\mathcal{H}$  is seen as a space of *private* messages : a private *dialect* useful only for  $\Pi^\bullet$  but that cannot be communicated to the environment ; the dialect is a typical additive creation, coming from the need of avoiding overlap in the treatment of the  $\&$ -rule.
- This distinction between these two components is a consequence of the *lamination* of  $\Pi^\bullet$  : by definition 8 every wire of  $\Pi^\bullet$  is of the form  $p(t, u) \mapsto q(t', u)$ .

In fact the privacy of the dialect is expressed by the fact that all constructions of geometry of interaction indeed deal with equivalence classes modulo  $\sim$ . In the case of binary rules, this fact limits the possible ways of merging dialects. If we want to combine  $U$  and  $V$ , we shall try to replace them by *variants*  $U'$  and  $V'$  whose respective dialects are in certain relation. We basically can

use the constructions of definition 7, and these two possibilities lead to the interpretations of the  $\otimes$ - and  $\&$ -rules, that we now sketch : let us assume that  $U$  and  $V$  are laminated wirings interpreting cut-free proofs of  $\vdash \Gamma, A$  and  $\vdash \Delta, B$  ; then

- The first operation is to create a common dialect ; this is done by :
  - in the case of a  $\otimes$ -rule, tensorization of the respective dialects : let  $U' = \otimes_1(U)$ ,  $V' = \otimes_2(V)$ .
  - in the case of a  $\&$ -rule, summation of the respective dialects :  $U' = \&_1(U)$ ,  $V' = \&_2(V)$ .
- The next operation is to glue together  $A$  and  $B$  ; this is done by merging (in  $U' + V'$ ) the predicates  $p_A$  and  $p_B$  respectively associated to  $A$  and  $B$  into a new predicate  $p_C$  (where  $C$  is  $A \otimes B$  or  $A \& B$ , depending on the rule), as follows :
  - every occurrence  $p_A(t, u)$  of  $p_A$  in  $U'$  is replaced with an occurrence of  $p_C(tg, u)$ ,
  - every occurrence  $p_B(t, u)$  of  $p_B$  in  $V'$  is replaced with an occurrence of  $p_C(td, u)$ .

the result  $W$  of this operation is the desired interpretation.

To end with our informal description of the interpretation, let us consider the case of a *CUT*-rule : let us assume that  $U$  and  $V$  interpret cut-free proofs of  $\vdash \Gamma, A$  and  $\vdash \Delta, A^\perp$  ; then

- We define  $U' = \otimes_1(U)$ ,  $V' = \otimes_2(V)$  as in the case of a  $\otimes$ -rule, and we sum them ;
- We introduce  $\sigma$  as the sum  $(p_A(x, y) \mapsto p_{A^\perp}(x, y)) + (p_{A^\perp}(x, y) \mapsto p_A(x, y))$ . The pair  $(U' + V', \sigma)$  is the desired interpretation.

### Remark 8

The rules for  $\otimes$  and  $\&$  share the last step, which consists in merging two predicates  $p_A$  and  $p_B$  by means of renaming of the channels ; the choice of the constants  $g$  and  $d$  to do so is arbitrary, but once it has been set, one cannot touch it. The main difference between a  $\otimes$ -rule and a  $\&$ -rule is the way in which the respective dialects are merged

- In the case of the  $\otimes$ , this is splendid ignorance (the dialects now commute with each other) : this is achieved by replacing  $u$  by  $uy$  and  $v$  by  $xv$  ;

- Whereas in the case of the  $\&$ , the dialects are just disjoint (i.e. incompatible). This is achieved by replacing  $u$  by  $ug$  and  $v$  by  $vd$ .

Observe that we could as well choose isomorphic solutions for the first step, for instance :

- Exchange (in the case of  $\otimes$ )  $\otimes_1$  with  $\otimes_2$ .
- Replace (in the case of  $\&$ )  $\&_1$  with  $\&_2$ .

This flexibility contrasts with the (relative) rigidity of the choice involved when merging predicates. This is because channels are public whereas dialects are private.

### 3.2 Interpretation of proofs of MALL

We first interpret the fragment **MALL** of multiplicative-additive linear logic.

#### Definition 9 (Pattern)

*The general pattern is as follows :*

- *We are given a proof  $\Pi$  of a sequent  $\vdash \Gamma$ , containing cuts on formulas  $\Delta$  ; in  $\Delta$  formulas are coupled in pairs  $(B, B^\perp)$  ; in fact we use the formulas of  $\Delta$  as labels for the cuts performed rather than actual formulas. The expression “Let  $\Pi$  be a proof of  $\vdash [\Delta], \Gamma$ ” is short for “Let  $\Pi$  be a proof of  $\vdash \Gamma$ , containing cuts on the formulas  $\Delta$ ”.*
- *We introduce the algebra  $\lambda^*(\Delta, \Gamma)$  :*
  - *the function letters are the constants  $d$  and  $g$  and the binary  $\odot$  ;*
  - *the predicate letters are all binary :  $p_A$  for all formula  $A$  in  $\Delta$  and  $\Gamma$ . We assume that these formulas are pairwise distinct (if not, they can be distinguished by adding indices).*
  - *an important particular case is  $\Delta = \emptyset$ , in which case we use the notation  $\lambda^*(\Gamma)$  ; if  $\Gamma$  were empty as well (a case that will never occur) then we would need to define  $\lambda^*(\emptyset)$ , that can conveniently be taken as the field  $\mathcal{C}$  of scalars.*
- *We define  $\sigma_{\Delta; \Gamma}$  in the algebra  $\lambda^*(\Delta, \Gamma)$  by*

$$\sigma_{\Delta; \Gamma} = \sum_{B \in \Delta} p_B(x, y) \mapsto p_{B^\perp}(x, y)$$

*$x, y$  are distinct variables, the sum is taken over all cut formulas ; observe that  $\sigma = \sigma^*$ , and that  $\sigma^2$  is a projection. Therefore  $1 - \sigma^2$  is a projection too, and*

$$\forall u \in \lambda^*(\Delta, \Gamma), \quad (1 - \sigma^2)u(1 - \sigma^2) \in \lambda^*(\Gamma)$$

in other terms prefixing and suffixing with  $1 - \sigma^2$  acts like a restriction, see subsection 5.7 in appendix.

- We interpret  $\Pi$  by a pair  $(\Pi^\bullet, \sigma_{\Delta; \Gamma})$  in the algebra  $\lambda^*(\Delta, \Gamma)$  : the definition of  $\Pi^\bullet$  is by induction.

**Definition 10 ( $\Pi^\bullet$ )**

Let  $\Pi$  be a proof of  $\vdash [\Delta], \Gamma$  ; we associate to  $\Pi$  its interpretation  $(\Pi^\bullet, \sigma_{\Delta; \Gamma})$ , where  $\Pi^\bullet$  is defined by induction :

**Case 10.1** If  $\Pi$  consists in the axiom  $\vdash A, A^\perp$ , then  $\Pi^\bullet = \sigma_{A, A^\perp}$ ; i.e.

$$\Pi^\bullet = (p_A(x, y) \mapsto p_{A^\perp}(x, y)) + (p_{A^\perp}(x, y) \mapsto p_A(x, y))$$

**Case 10.2** If  $\Pi$  is obtained from  $\Pi_1$  and  $\Pi_2$  by means of a CUT-rule with cut-formulas  $A, A^\perp$ , then

- let us replace in  $\Pi_1^\bullet$  all atoms  $p_C(t, u)$  by  $p_C(t, uy)$  ( $y$  fresh variable) ; the result is called  $U_1$  ;
- let us replace in  $\Pi_2^\bullet$  all atoms  $p_C(t, v)$  by  $p_C(t, xv)$  ( $x$  fresh variable) ; the result is called  $U_2$  ; then

$$\Pi^\bullet = U_1 + U_2.$$

Of course the essential feature of this step is the modification of the component  $\sigma$  : if  $\Pi_1$  and  $\Pi_2$  respectively prove  $\vdash [\Delta_1], \Gamma_1, A$  and  $\vdash [\Delta_2], \Gamma_2, A^\perp$ , then  $\Pi$  is a proof of  $\vdash [\Delta_1, \Delta_2, A, A^\perp], \Gamma_1, \Gamma_2$ .

**Case 10.3** If  $\Pi$  is obtained from  $\Pi_1$  by means of an exchange rule, then

$$\Pi^\bullet = \Pi_1^\bullet.$$

**Case 10.4** If  $\Pi$  is obtained from  $\Pi_1$  and  $\Pi_2$  by means of a  $\otimes$ -rule introducing  $A \otimes B$ , then

- let us replace in  $\Pi_1^\bullet$  all atoms  $p_C(t, u)$  by  $p_C(t, uy)$  ( $y$  fresh variable) when  $C$  is distinct from  $A$ , and all atoms  $p_A(t, u)$  by  $p_{A \otimes B}(tg, uy)$  ; the result is called  $U_1$  ;
- let us replace in  $\Pi_2^\bullet$  all atoms  $p_C(t, v)$  by  $p_C(t, xv)$  ( $x$  fresh variable) when  $C$  is distinct from  $B$ , and all atoms  $p_B(t, v)$  by  $p_{A \otimes B}(td, xv)$  ; the result is called  $U_2$  ; then

$$\Pi^\bullet = U_1 + U_2.$$

**Case 10.5** If  $\Pi$  is obtained from  $\Pi_1$  by means of a  $\wp$ -rule, then let us replace in  $\Pi_1^\bullet$  all atoms  $p_A(t, u)$  by  $p_{A\wp B}(tg, u)$  and all atoms  $p_B(t, v)$  by  $p_{A\wp B}(td, v)$  ; the result is by definition  $\Pi^\bullet$ .

**Case 10.6** If  $\Pi$  is obtained from  $\Pi_1$  and  $\Pi_2$  by means of a  $\&$ -rule introducing  $A\&B$ , then

- let us replace in  $\Pi_1^\bullet$  all atoms  $p_C(t, u)$  by  $p_C(t, ug)$  when  $C$  is distinct from  $A$ , and all atoms  $p_A(t, u)$  by  $p_{A\&B}(tg, ug)$  ; the result is called  $U_1$  ;
- let us replace in  $\Pi_2^\bullet$  all atoms  $p_C(t, v)$  by  $p_C(t, vd)$  when  $C$  is distinct from  $B$ , and all atoms  $p_B(t, v)$  by  $p_{A\&B}(td, vd)$  ; the result is called  $U_2$  ; then

$$\Pi^\bullet = U_1 + U_2.$$

**Case 10.7** • If  $\Pi$  is obtained from  $\Pi_1$  by means of a  $\oplus_1$ -rule, then let us replace in  $\Pi_1^\bullet$  all atoms  $p_A(t, u)$  by  $p_{A\oplus B}(tg, u)$  ; the result is by definition  $\Pi^\bullet$ .

- If  $\Pi$  is obtained from  $\Pi_1$  by means of a  $\oplus_2$ -rule, then let us replace in  $\Pi_1^\bullet$  all atoms  $p_B(t, u)$  by  $p_{A\oplus B}(td, u)$  ; the result is by definition  $\Pi^\bullet$ .

**Case 10.8** If  $\Pi$  is obtained from  $\Pi'$  by means of  $\forall$ -rule (first or second order), then  $\Pi^\bullet = \Pi'^\bullet$ .

**Case 10.9** If  $\Pi$  is obtained from  $\Pi'$  by means of  $\exists$ -rule (first or second order), then  $\Pi^\bullet = \Pi'^\bullet$ .

### Remark 9

Definition 10 uses several times rather arbitrary constructions to merge dialects (the cases 10.4 and 10.6, but also the constant  $g$  in the case 10.1) . In fact in all these cases, we might have chosen variants, without any difference. This is because of proposition 1 will allows one to replace  $\Pi^\bullet$  with a variant.

### Example 1

Cut-elimination replaces the cut

$$\frac{\frac{}{\vdash \Gamma, A} \quad \frac{}{\vdash B^\perp, B} Id}{\vdash \Gamma, B} CUT$$

(with  $A$  and  $B$  different occurrences of the same formula) by the original proof  $\Pi$  of  $\vdash \Gamma, A$ . The original proof is interpreted as

$(\otimes_1(\Pi^\bullet) + \otimes_2(\sigma_{B^\perp, B} ; \Gamma, A), \sigma_{A, B^\perp} ; \Gamma, B)$ . Nilpotency is easily shown  $((U\sigma)^2 = 0)$ , and the result of execution is  $(\otimes_1(\Pi^\bullet))$  in which  $p_A$  has been replaced with  $p_B$ , i.e. it is a variant of the interpretation of the modified proof. This shows, in this basic, but essential example, that execution corresponds to cut-elimination, up to variance.



### 3.3 The nilpotency theorem

Our first goal is to prove that  $EX(\Pi^\bullet, \sigma)$  ( $= \Pi^\bullet(\sigma\Pi^\bullet)^{-1}$ , see definition 5) makes sense, i.e.,  $\Pi^\bullet\sigma$  is *nilpotent*. This will be done by an adaptation of the results from [G88].

#### Theorem 1 (Transitivity of cut)

In the algebra  $\lambda^*(\Delta, \Delta', \Gamma)$  let us define  $\sigma, \tau$  as respectively  $\sigma_{\Delta; \Delta', \Gamma}$  and  $\sigma_{\Delta'; \Delta, \Gamma}$  (so that  $\sigma_{\Delta, \Delta'; \Gamma} = \sigma + \tau$ ) ; let  $U$  be a wiring

- If  $U\sigma$  is nilpotent,  $RES(U, \sigma)$  ( $= (1 - \sigma^2)EX(U, \sigma)(1 - \sigma^2)$ ) is a wiring.
- $(\sigma + \tau)U$  is nilpotent iff  $\sigma U$  and  $\tau \cdot RES(U, \sigma)$  are nilpotent ;
- In that case  $RES(U, \sigma + \tau) = RES(RES(U, \sigma), \tau)$ .

PROOF: see [G88], lemmas 4 and 5, where this result was called *associativity of cut*.  $\square$

#### Theorem 2 (Nilpotency)

$\sigma\Pi^\bullet$  is nilpotent.

PROOF: the proof is adapted from [G88]. (We neglect all features related to quantifiers, which are quite the same.) We first define :

#### Definition 11 (Weak orthogonality)

Let  $U$  and  $V$  be laminated wirings in  $\lambda^*(T^2)$  (an algebra with a single binary predicate  $p$ ;  $U$  and  $V$  are said to be weakly orthogonal if  $\otimes_1(U) \cdot \otimes_2(V)$  is nilpotent. In that case we introduce the notation  $U \perp V$ .

The definition makes sense because of the following :

#### Proposition 3

- If  $U \perp V$ , then  $V \perp U$ .
- If  $U \sim U'$  and  $V \sim V'$  and  $U \perp V$ , then  $U' \perp V'$ .

#### Definition 12 (Weak types)

Given a subset  $X$  of  $\lambda^*(T^2)$ , we define  $X^\perp = \{U ; \forall V \in X, U \perp V\}$ .

A weak type is any set  $X$  of laminated wirings such that  $X = X^{\perp\perp}$ . The connectives of linear logic can be interpreted as operations on weak types :

**Case 12.1** If  $X$  and  $Y$  are weak types, then we define the weak type  $X \otimes Y$  as  $Z^{\perp\perp}$ , where  $Z$  consists of all  $U' + V'$ , where  $U'$  and  $V'$  are constructed from  $U \in X$  and  $V \in Y$  as follows :

- Replace in  $U$  all atoms  $p(t, u)$  by  $p(tg, uy)$ .
- Replace in  $U$  all atoms  $p(t, v)$  by  $p(td, xv)$ .

**Case 12.2** If  $X$  and  $Y$  are weak types, then we define the weak type  $X \wp Y$  as  $(X^\perp \otimes Y^\perp)^\perp$ .

**Case 12.3** If  $X$  and  $Y$  are weak types, then we define the weak type  $X \oplus Y$  as  $Z^{\perp\perp}$ , where  $Z$  consists of : all  $U'$  and all  $V'$ , where  $U'$  and  $V'$  are constructed from  $U \in X$  and  $V \in Y$  as follows :

- Replace in  $U$  all atoms  $p(t, u)$  by  $p(tg, u)$ .
- Replace in  $U$  all atoms  $p(t, v)$  by  $p(td, v)$ .

**Case 12.4** If  $X$  and  $Y$  are weak types, then we define the weak type  $X \& Y$  as  $(X^\perp \oplus Y^\perp)^\perp$ .

### Remark 10

The interpretation by weak types is very degenerated ; in particular, it easy to check that  $\&$  and  $\oplus$  are not separated by this rough interpretation.

### Definition 13 (Weak types associated to formulas and sequents)

Let us associate an arbitrary weak type  $\alpha^\bullet$  to each atom  $\alpha$  of a language in linear logic ; then each formula  $A$  built from such atoms by means of linear negation, and the binary connectives  $\otimes, \wp, \&, \oplus$  immediately gets interpreted by a weak type  $A^\bullet$ . (This also extends to quantifiers of any order, provided one works on weak type parameters as in definition 3 of [G88].) The definition of  $A^\bullet$  immediately induces a definition of  $\Gamma^\bullet$ , when  $\vdash \Gamma$  is a sequent : assume that  $\Gamma = A_1, \dots, A_n$ , and let  $U$  be a laminated wiring in  $\lambda^*(\Gamma)$  ;

- Given laminated wirings  $V_1, \dots, V_n$  in  $A_1^{\perp\bullet}, \dots, A_n^{\perp\bullet}$  respectively, we can rename the predicates so that  $V_i \in \lambda^*(A_i)$ , and introduce  $V'_1 + \dots + V'_n$  as the result of applying  $n - 1$  cases 10.4 to the  $V_i$  ; the object depends in which order the rules are performed, but different choices yield variants.
- Then  $U \in \Gamma^\bullet$  iff for all  $V_1, \dots, V_n$  in  $A_1^{\perp\bullet}, \dots, A_n^{\perp\bullet}$  respectively,  $U \perp V'_1 + \dots + V'_n$ .

of course the definition has two particular cases of interest :

- The particular case  $n = 1$  yields  $(\vdash A)^\bullet = A^\bullet$  ;
- The particular case  $n = 0$  yields  $\vdash^\bullet = \mathcal{C}$ , see definition 9.

**Lemma 2.1** *If  $U, V$  are variants and  $U \in A^\bullet$  (resp.  $U \in \Gamma^\bullet$ ), then  $V \in A^\bullet$  (resp.  $V \in \Gamma^\bullet$ ).*

PROOF: immediate  $\square$

**Lemma 2.2**  *$U \in (\Gamma, A)^\bullet$  iff for all  $V \in A^{\perp\bullet}$  the “cut”  $(U_1 + V_1, \sigma_{A, A^\perp; \Gamma})$  (defined as in case 10.2) is such that :  $(U_1 + V_1)\sigma_{A, A^\perp; \Gamma}$  is nilpotent and the result  $RES(U_1 + V_1, \sigma_{A, A^\perp; \Gamma})$  is in  $\Gamma^\bullet$ .*

PROOF: easy, see lemma 7 of [G88].  $\square$

The nilpotency theorem follows from the more precise :

### Theorem 3

*If  $\Pi$  is a proof of  $\vdash [\Delta], \Gamma$ , then  $\Pi^\bullet \sigma_{\Delta; \Gamma}$  is nilpotent and  $RES(\Pi^\bullet, \sigma_{\Delta; \Gamma}) \in \Gamma^\bullet$ .*

PROOF: this theorem is proved by induction on  $\Pi$  ; there is a general pattern that can be followed for all cases. We do this in detail for the first case (the case of a  $\&$ -rule, which is the truly new case), but we shall only treat simplified versions of the other cases. The pattern basically reduces the general case to the cut-free case (i.e. when the premise(s) is (are) cut-free), and then the cut-free case to a context-free case. To do this in very rigorous way, it would be natural to enlarge the set of rules with an additional one : for each  $V \in \vdash \Gamma^\bullet$  add the axiom  $\Gamma$  whose proof  $\Pi_V$  is interpreted as  $\Pi_V^\bullet = V$  (in the spirit of model-theory, where new constants for the elements of the model are introduced).

**Case 3.1** • We first treat a very limited case : the last rule of  $\Pi$  is a  $\&$ -rule between proofs  $\Pi_1$  of  $\vdash A_1$  and  $\Pi_2$  of  $\vdash A_2$ , (so that  $\Delta = \emptyset$  and  $\Gamma = A_1 \& A_2$ ). We have to show that for all  $U'$  obtained from  $U \in A_1^{\perp\bullet}$  (resp. all  $V'$  obtained from  $V \in A_2^{\perp\bullet}$ ) by means of definition 12.3, then  $\Pi^\bullet \perp U'$  (resp.  $\Pi^\bullet \perp V'$ ). But this immediately reduces to  $\Pi_1^\bullet \perp U$  (resp.  $\Pi_2^\bullet \perp V$ ).

- Then we treat the case where the premises are cut-free : the last rule of  $\Pi$  is a  $\&$ -rule between proofs  $\Pi_1$  of  $\vdash \Phi, A_1$  and  $\Pi_2$  of  $\vdash \Phi, A_2$ , (so that  $\Delta = \emptyset$  and  $\Gamma = \Phi, A_1 \& A_2$ ). We can argue by induction on the cardinality of  $\Phi$  :

- the case  $\Phi = \emptyset$  has already been treated ;
- if  $\Phi = B, \Psi$ , then by lemma 2.2,  $\Pi^\bullet \in (\Phi, A_1 \& A_2)^\bullet$  (resp.  $\Pi_1^\bullet \in (\Phi, A_1)^\bullet, \Pi_2^\bullet \in (\Phi, A_2)^\bullet$ ) iff for all  $V \in B^{\perp\bullet}$   $\Pi^\bullet \sigma_{B, B^\perp; \Psi, A_1 \& A_2}$  is nilpotent and  $RES(\Pi^\bullet, \sigma_{B, B^\perp; \Psi, A_1 \& A_2}) \in (\Psi, A_1 \& A_2)^\bullet$  (resp.  $\Pi_1^\bullet \sigma_{B, B^\perp; \Psi, A_1}$  is nilpotent and  $RES(\Pi_1^\bullet, \sigma_{B, B^\perp; \Psi, A_1}) \in (\Psi, A_1)^\bullet$ ,  $\Pi_2^\bullet \sigma_{B, B^\perp; \Psi, A_2}$  is nilpotent and  $RES(\Pi_2^\bullet, \sigma_{B, B^\perp; \Psi, A_2}) \in (\Psi, A_2)^\bullet$ ). The result follows from the induction hypothesis on the size of  $\Phi$  and the fact that

- \*  $\Pi^\bullet \sigma_{B,B^\perp;\Psi,A_1 \& A_2}$  is nilpotent iff  $\Pi_1^\bullet \sigma_{B,B^\perp;\Psi,A_1}$  and  $\Pi_2^\bullet \sigma_{B,B^\perp;\Psi,A_2}$  are nilpotent ;
- \* if we apply definition 10.6 to  $RES(\Pi_1^\bullet, \sigma_{B,B^\perp;\Psi,A_1})$  and  $RES(\Pi_2^\bullet, \sigma_{B,B^\perp;\Psi,A_2})$ , then we get a variant of  $RES(\Pi^\bullet, \sigma_{B,B^\perp;\Psi,A_1 \& A_2})$ .

- It remains to treat the general case : the last rule of  $\Pi$  is a  $\&$ -rule between proofs  $\Pi_1$  of  $\vdash [\Delta_1], \Phi, A_1$  and  $\Pi_2$  of  $\vdash [\Delta_2], \Phi, A_2$ , (so that  $\Delta = \Delta_1, \Delta_2$  and  $\Gamma = \Phi, A_1 \& A_2$ ). From the induction hypothesis,  $\Pi_i^\bullet \sigma_{\Delta_i, \Phi, A_i}$  is nilpotent and  $RES(\Pi_i^\bullet, \sigma_{\Delta_i, \Phi, A_i}) \in (\Phi, A_i)^\bullet$  for  $i = 1, 2$ . Let  $u_i = RES(\Pi_i^\bullet, \sigma_{\Delta_i, \Phi, A_i})$  ; then it is easily seen that  $\Pi^\bullet \sigma_{\Delta, \Phi, A \& B}$  is nilpotent and that by definition 10.6 applied to  $u_1$  and  $u_2$ , we get  $RES(\Pi^\bullet, \sigma_{\Delta, \Phi, A \& B}) \in (\Phi, A \& B)^\bullet$ .

**Case 3.2** If  $\Pi$  is the axiom  $\vdash A, A^\perp$  ; for reasons of labelling, we prefer to rename the first  $A$  as  $B$ ; then

$$\Pi^\bullet = (p_B(x, y) \mapsto p_{A^\perp}(x, y)) + (p_{A^\perp}(x, y) \mapsto p_B(x, y))$$

select  $U = \sum p_A(t_i, v_i) \mapsto p_A(u_i, v_i)$  in  $A^\bullet$  ; then these two wirings are respectively modified into

$$\otimes_1(\Pi^\bullet) = (p_B(x, yz) \mapsto p_{A^\perp}(x, yz)) + (p_{A^\perp}(x, yz) \mapsto p_B(x, yz))$$

and  $\otimes_2(U) = \sum p_A(t_i, y'v_i) \mapsto p_A(u_i, y'v_i)$ . The nilpotency of  $(\otimes_1(\Pi^\bullet) + \otimes_2(U))\sigma_{A, A^\perp; B}$  is rather immediate, and the result is easily seen to be  $\sum p_B(t_i, yv_i) \mapsto p_B(u_i, yv_i)$ , a variant of  $U$  ; the result follows from lemmas 2.2 and 2.1.

**Case 3.3** The last rule of  $\Pi$  is a *CUT*-rule between proofs  $\Pi_1$  of  $\vdash A$  and  $\Pi_2$  of  $\vdash A^\perp$ . (This extremely simplified case can only occur because we have extended the set of all proofs by adding a lot of new axioms ; also remark that the reduction of the general case to this case makes a heavy use of transitivity of cut). Due to the peculiar definition of  $\vdash^\bullet$ , it suffices to show that  $\Pi^\bullet \sigma_{A, A^\perp}$  is nilpotent ; but this immediately reduces to the orthogonality of  $\Pi_1^\bullet$  and  $\Pi_2^\bullet$ , which is precisely the induction hypothesis.

**Case 3.4** The last rule of  $\Pi$  is a  $\oplus_1$ -rule applied to a proof  $\Pi_1$  of  $\vdash A$ . This simplified case follows immediately from definition 12.3. The case of a  $\oplus_2$ -rule is similar.

**Case 3.5** The last rule of  $\Pi$  is a  $\otimes$ -rule between proofs  $\Pi_1$  of  $\vdash A$  and  $\Pi_2$  of  $\vdash B$ . This simplified case follows from definition 12.1.

**Case 3.6** The last rule of  $\Pi$  is a  $\wp$ -rule applied to a proof  $\Pi_1$  of  $\vdash A, B$ . By definition 12.2 we must show that for all  $U \in A^{\perp\bullet}$  and  $V \in B^{\perp\bullet}$ , then (with  $U'$  and  $V'$  as in definition 12.3)  $\Pi^\bullet \perp U' + V'$ . But this is easily reduced to  $\Pi_1^\bullet \perp \otimes_1(U) + \otimes_2(V)$ , i.e. the induction hypothesis.

**Case 3.7** The remaining cases are left to the reader.  $\square$

### Remark 11

Although the result proved (nilpotency) is rather weak, it is “proof-theoretically strong” : for instance, in the presence of exponentials (next subsection) and second order quantification, it cannot be proved inside second-order arithmetic : this is because the orders of nilpotency are related to the length of normalisation steps (see [DR93] for instance).

## 3.4 $\flat$ -sequent calculus

Our next goal is to prove that the result  $RES(\Pi^\bullet, \sigma)$  of the execution is  $\Pi_1^\bullet$ , where  $\Pi_1$  is obtained from  $\Pi$  by cut-elimination. Unfortunately

- We can only expect a *variant* of  $\Pi_1^\bullet$  ;
- Worse, the result  $RES(\Pi^\bullet, \sigma)$  contains “un erased” information, what we called a *beard* in [G88]. This second feature makes the precise statement of a theorem quite delicate.

Since of the two partners -geometry of interaction and cut-elimination-, the former (although very recent) is the most natural, we shall therefore (slightly) modify sequent calculus : the result is called  $\flat$ -sequent calculus. The reason for this terminology is that the calculus contains unessential rules corresponding to the erasings that geometry of interaction, in its lazy spirit, does not perform. We introduce a new symbol,  $\flat$ , that we call *flat* ; this symbol is treated differently from the other formulas of linear logic :  $\flat$  cannot be combined ( $\flat \wp \flat$ ,  $\flat^\perp$  are not allowed) ; this leaves only the possibility of using  $\flat$  inside sequents, like in  $\vdash \flat$ ,  $\vdash \flat, A$ ,  $\vdash \flat, A, \flat$  etc. The specific rules of  $\flat$  are the following :

$$\frac{}{\vdash \flat, \Gamma} \flat \quad \frac{\vdash \flat, \flat, \Gamma}{\vdash \flat, \Gamma} c\flat \quad \frac{\vdash \flat, \Gamma \quad \vdash \Gamma}{\vdash \Gamma} s\flat$$

- The axiom scheme  $\vdash \flat, \Gamma$  for any  $\Gamma$ .
- The contraction rule : from  $\vdash \flat, \flat, \Gamma$  deduce  $\vdash \flat, \Gamma$ .
- The summation rule : from  $\vdash \flat, \Gamma$  and  $\vdash \Gamma$  deduce  $\vdash \Gamma$ .

These rules are basically minor structural rules.

**Remark 12**

- $\flat$  can be faithfully interpreted as the constant  $\top$  of linear logic : the axiom is the usual axiom for  $\top$  and the rules are derivable... but this is not the point.
- Among the derivable rules for  $\flat$  :
  - a form of inflation : from  $\vdash \flat, \Gamma$  derive  $\vdash \flat, \Gamma, \Delta$
  - another form of summation : from  $\vdash \flat, \Gamma$  and  $\vdash \flat, \Gamma$  deduce  $\vdash \flat, \Gamma$ .
- Contraction and inflation imply that we do not count the multiplicity of  $\flat$ . This means that in reality, we are working with two kinds of sequents : usual ones (without  $\flat$ ) and *flat* ones (with at least one  $\flat$ ). Of course we get a much simpler presentation with only one kind of sequents, and this explains our choice.

**Definition 14 (Geometry of interaction of  $\flat$ )**

A proof  $\Pi$  of  $\vdash [\Delta], \Gamma$  is interpreted as a laminated wire in  $\lambda^*(\Delta, \Gamma - \flat)$ , where  $\Gamma - \flat$  is obtained from  $\Gamma$  by the removal of all  $\flat$  (observe that since  $\flat$  cannot be negated, it does not occur in  $\Delta$ ).

- An axiom  $\vdash \flat, \Gamma$  is interpreted by  $0 \in \lambda^*(\Gamma - \flat)$ .
- Contraction is interpreted identically.
- If  $\Pi$  follows from  $\Pi_1$  and  $\Pi_2$  by means of a summation : then  $\Pi^\bullet = \&_1(\Pi_1^\bullet) + \&_2(\Pi_2^\bullet)$ .

**Remark 13**

The basic idea behind the symbol “ $\flat$ ” is that a proof of a flat sequent  $\vdash \flat, \Gamma$  comes from a “real” proof in which some essential part has been erased. Flat proofs interact with real ones by means of summation : summation typically occurs in the case of cut-elimination on an additive formula. For instance assume that we perform a cut between a proof of  $\vdash A \oplus B$  (obtained from a proof  $\Pi$  of  $A$  by a  $\oplus_1$ -rule) and a proof of of a sequent  $\vdash A^\perp \& B^\perp, \Gamma$  coming from proofs  $\Pi_1$  of  $\vdash A^\perp, \Gamma$  and  $\Pi_2$  of  $\vdash B^\perp, \Gamma$ . The traditional way of eliminating the cut is to replace it with the proof  $\Pi'$  obtained by a cut between  $\Pi$  and  $\Pi_1$ , but this does not quite correspond to geometry of interaction. This is why in the  $\flat$ -calculus, our cut is replaced by the *summation* of  $\Pi'$  and the proof of  $\vdash \flat, \Gamma$  obtained by a cut between the axiom  $\vdash \flat, B$  and  $\Pi_2$ .

It is not surprising that the  $\mathfrak{b}$ -calculus enjoys cut-elimination. . . In fact we are rather interested in the details of the cut-elimination procedure that we now sketch :

**Definition 15 (Cut-elimination procedure)**

A cut

$$\frac{\frac{\cdots}{\vdash \Gamma, A} R \quad \frac{\cdots}{\vdash \Delta, A^\perp} S}{\vdash \Gamma, \Delta} CUT$$

is replaced as follows :

- If the rule  $R$  is not an introduction of  $A$  or an axiom, then  $R$  is performed after the cut, for instance :

- If  $R$  is a summation rule, with premises  $\vdash \Gamma, A$  and  $\vdash \Gamma, A, \mathfrak{b}$ , then our cut is replaced with :

$$\frac{\frac{\vdash \Gamma, A \quad \vdash \Delta, A^\perp}{\vdash \Gamma, \Delta} CUT \quad \frac{\vdash \Gamma, A, \mathfrak{b} \quad \vdash \Delta, A^\perp}{\vdash \Gamma, \Delta, \mathfrak{b}} CUT}{\vdash \Gamma, \Delta} s\mathfrak{b}$$

- If the rule  $S$  is not an introduction of  $A^\perp$  or an axiom, then  $S$  is performed after the cut, as above.
- If both  $R$  and  $S$  are introductions of the cut-formulas or axioms, then there are several cases, including :

- when  $R$  is the identity axiom  $\vdash A, A^\perp$ , the proof is replaced with  $\frac{}{\vdash \Delta, A^\perp} S$  ;
- when  $R$  is the axiom  $\vdash \Phi, A, \mathfrak{b}$ , the solution depends on  $S$ , typically :
  - \* if  $S$  is the axiom  $\vdash \Psi, A^\perp, \mathfrak{b}$ , then we replace the proof with the axiom  $\vdash \Phi, \Psi, \mathfrak{b}, \mathfrak{b}$  ;
  - \* if  $S$  is a  $\otimes$ -rule, with premises  $\vdash \Psi, B^\perp$  and  $\vdash \Xi, C^\perp$ , then our proof is replaced with two cuts between the axiom  $\vdash \Phi, B, C, \mathfrak{b}$  and the premises of the  $\otimes$ -rule ;
  - \* if  $S$  is a  $\wp$ -rule, with premise  $\vdash \Delta, B^\perp, C^\perp$ , then our proof is replaced with two cuts between the axioms  $\vdash \Phi, B, \mathfrak{b}$  and  $\vdash C, \mathfrak{b}$  and the premise of the  $\wp$ -rule ; this yields a proof of  $\vdash \Phi, \Delta, \mathfrak{b}, \mathfrak{b}$  that we contract into  $\vdash \Phi, \Delta, \mathfrak{b}$  ;

- \* if  $S$  is a  $\&$ -rule, with premises  $\vdash \Delta, B^\perp$  and  $\vdash \Delta, C^\perp$ , then we cut these premises respectively with the axioms  $\vdash \Phi, B^\perp, \flat$  and  $\vdash \Phi, C^\perp, \flat$ , yielding two proofs of  $\vdash \Phi, \Delta, \flat$  to which we can apply summation, as seen in remark 12 ;
- \* if  $S$  is a  $\oplus_1$ -rule, with premise  $\vdash \Delta, B^\perp$ , then our proof is replaced with a cut between this premise and the axiom  $\vdash \Phi, B, \flat \dots$
- when  $R$  and  $S$  are both logical rules, then we get several cases, including :

- \* if  $R$  is a  $\wp$ -rule, with premise  $\vdash \Gamma, B, C$  and  $S$  is a  $\otimes$ -rule, with premises  $\vdash \Psi, B^\perp$  and  $\vdash \Xi, C^\perp$ , then we replace our proof with

$$\frac{\frac{\vdash \Gamma, B, C \quad \vdash \Psi, B^\perp}{\vdash \Gamma, \Psi, C} CUT \quad \vdash \Xi, C^\perp}{\vdash \Gamma, \Psi, \Xi} CUT$$

- \* if  $R$  is a  $\oplus_1$ -rule, with premise  $\vdash \Gamma, B$  and  $S$  is a  $\&$ -rule, with premises  $\vdash \Delta, B^\perp$  and  $\vdash \Delta, C^\perp$ , then we replace our proof with

$$\frac{\frac{\vdash \Gamma, B \quad \vdash \Delta, B^\perp}{\vdash \Gamma, \Delta} CUT \quad \frac{\frac{\vdash \Gamma, C, \flat}{\vdash \Gamma, \Delta, \flat} CUT \quad \frac{\vdash \Delta, C^\perp}{\vdash \Gamma, \Delta, \flat} CUT}{\vdash \Gamma, \Delta} sb$$

#### Proposition 4

*Cut-elimination holds for the  $\flat$ -calculus, using the procedure sketched in definition 15.*

PROOF: boring and straightforward ; however notice that the treatment of a cut on an additive formula makes a significant use of flat sequents.  $\square$

#### Theorem 4 (Soundness)

*If a cut-free proof  $\Pi$  of  $\vdash \Gamma$  is obtained from a proof  $\Sigma$  of  $\vdash [\Delta], \Gamma$  by means of the transformations sketched in definition 15, then  $\Pi^\bullet$  is a variant of  $RES(\Sigma^\bullet, \sigma_{\Delta; \Gamma})$ .*

**Lemma 4.1** *If a proof  $\Pi$  of  $\vdash [\Delta'], \Gamma$  is obtained from a proof  $\Sigma$  of  $\vdash [\Delta], \Gamma$  by one step of definition 15, then  $RES(\Pi^\bullet, \sigma_{\Delta'; \Gamma}) \sim RES(\Sigma^\bullet, \sigma_{\Delta; \Gamma})$ .*

PROOF: we treat only a few distinguished cases :



- Assume that  $\Sigma$  is obtained from  $\Sigma_1, \Sigma_2$  respectively proving  $\vdash [\Delta], \Phi, A$ , and  $\vdash A^\perp, A$ , (an identity axiom that we note  $\vdash A^\perp, B$ ) via a cut on  $A$ . Then  $\Sigma^\bullet$  is of the form  $\otimes_1(\Sigma_1^\bullet) + \otimes_2(\Sigma_2^\bullet)$ . Transitivity of cut yields  $RES(\Sigma^\bullet, \sigma_{\Delta, A, A^\perp; \Phi, B}) = RES(RES(\Sigma^\bullet, \sigma_{\Delta; A, A^\perp, \Phi, B}), \sigma_{A, A^\perp; \Delta, \Phi, B})$ . The result follows from the fact that  $RES(\otimes_1(U) + \otimes_2(\Sigma_2^\bullet), \sigma_{A, A^\perp; \Delta, \Phi, B})$  is a variant of  $U$  (with  $U = RES(\Sigma_1, \sigma_{\Delta; \Phi, A})$ ).
- Assume that  $\Sigma$  is obtained from  $\Sigma_1, \Sigma_2, \Sigma_3$ , respectively proving  $\vdash [\Delta_1], \Phi, B, C, \vdash [\Delta_2], \Psi, B^\perp$  and  $\vdash [\Delta_3], \Xi, C^\perp$ , by a  $\wp$ -rule and a  $\otimes$ -rule followed by a cut, then the construction of  $\Sigma^\bullet$  involves two steps :
  - the formation of the sum  $U = \otimes_1(\Sigma_1^\bullet) + \otimes_2(\otimes_1(\Sigma_2^\bullet) + \otimes_2(\Sigma_3^\bullet))$  ; this sum is a variant of  $\Pi^\bullet = \otimes_1(\otimes_1(\Sigma_1^\bullet) + \otimes_2(\Sigma_2^\bullet)) + \otimes_2(\Sigma_3^\bullet)$
  - the merge in  $U$  of  $p_B$  and  $p_C$  into  $p_A$  and of  $p_{B^\perp}$  and  $p_{C^\perp}$  into  $p_{A^\perp}$ .

If we define  $\Lambda = \Delta_1, \Delta_2, \Delta_3$ , then it is immediate that

$RES(\Sigma^\bullet, \sigma_{\Lambda, A, A^\perp; \Gamma}) = RES(U, \sigma_{\Lambda, B, B^\perp, C, C^\perp; \Gamma})$ . But  $\Lambda, B, B^\perp, C, C^\perp$  is  $\Delta'$  so  $RES(U, \sigma_{\Lambda, B, B^\perp, C, C^\perp; \Gamma})$  is a variant of  $RES(\Pi^\bullet, \sigma_{\Delta'; \Gamma})$ .

- Assume that  $\Sigma$  is obtained from  $\Sigma_1, \Sigma_2, \Sigma_3$ , respectively proving  $\vdash [\Delta_1], \Phi, B, C, \vdash [\Delta_2], \Psi, B^\perp$  and  $\vdash [\Delta_3], \Psi, C^\perp$ , by a  $\oplus_1$ -rule and a  $\&$ -rule followed by a cut, then the construction of  $\Sigma^\bullet$  involves two steps :
  - the formation of the sum  $U = \otimes_1(\Sigma_1^\bullet) + \otimes_2(\&_1(\Sigma_2^\bullet) + \&_2(\Sigma_3^\bullet))$  ; this sum is a variant of  $\Pi^\bullet = \&_1(\otimes_1(\Sigma_1^\bullet) + \otimes_2(\Sigma_2^\bullet)) + \&_2(\otimes_2(\Sigma_3^\bullet))$
  - the merge in  $U$  of  $p_B$  and  $p_C$  into  $p_A$  and of  $p_{B^\perp}$  and  $p_{C^\perp}$  into  $p_{A^\perp}$ .

If we define  $\Lambda = \Delta_1, \Delta_2, \Delta_3$ , then it is immediate that

$RES(\Sigma^\bullet, \sigma_{\Lambda, A, A^\perp; \Gamma}) = RES(U, \sigma_{\Lambda, B, B^\perp, C, C^\perp; \Gamma})$ . But  $\Lambda, B, B^\perp, C, C^\perp$  is  $\Delta'$  so  $RES(U, \sigma_{\Lambda, B, B^\perp, C, C^\perp; \Gamma})$  is a variant of  $RES(\Pi^\bullet, \sigma_{\Delta'; \Gamma})$ .

- All the petty cuts between an axiom for  $\flat$  and a logical rule (or an axiom for  $\flat$ ) are easily treated : this is because the axiom for flat is interpreted by 0.
- The endless list of commutations of the *CUT*-rule is easily handled : such steps introduce variants.  $\square$

PROOF: the proof of theorem 4 is immediate from the lemma.  $\square$

Whether or not we have totally achieved our task, surely execution corresponds exactly to normalisation, but in a slightly exotic sequent calculus. This calculus contains more cut-free proofs than usual, and therefore its use might be problematic. But the question is solved by the following proposition :

### Proposition 5

*In the  $\flat$ -calculus, the booleans remain standard : there is a sequent  $\Gamma$  such that the set  $\Gamma^\bullet$  of all  $\Pi^\bullet$  where  $\Pi$  varies through cut-free proofs of  $\Gamma$  has exactly two elements.*

PROOF: we have to give a precise meaning to the proposition.

- We can decide to represent booleans by a sequent  $\vdash \alpha^\perp, \alpha \oplus \alpha$ , where  $\alpha$  is a given atomic formula.
- This formula has only two cut-free proofs in the usual sequent calculus (Identity axiom and a  $\oplus_1$ -rule, identity axiom and a  $\oplus_2$ -rule) ; these two proofs can be taken as the two booleans. Geometry of interaction obviously interprets them differently.
- In the  $\flat$ -calculus, the summation rule offers more possibilities of cut-free proofs ; however it is easily proved that the flat summands must have interpretation 0.  $\square$

Therefore the execution formula is consistent with usual syntactic manipulations : this means that for any boolean question that we can ask of the output of some cut-elimination, both methods will yield the same answer. See [G88] for a discussion.

### 3.5 The neutral elements

The case of the multiplicative units is delicate ; the only “natural” choice for interpreting the rule of introduction of  $\perp$  is the trivial one : if a proof  $\Pi$  of  $\vdash \Gamma, \perp$  comes from a proof  $\vdash \Pi_1$  of  $\Gamma$  by the  $\perp$ -rule, let  $\Pi^\bullet = \Pi_1^\bullet$ . This choice has the following consequence : given two proofs  $\Pi_1$  and  $\Pi_2$  of  $\vdash \Gamma$  (this sequent is written below as  $\vdash \Gamma_i$  to distinguish the two proofs), the proofs

$$\frac{\frac{\vdash \Gamma_1}{\vdash \Gamma, \perp} \perp \quad \frac{\vdash \Gamma_2}{\vdash \Gamma, \perp} \perp}{\vdash \Gamma, \perp \& \perp} \& \quad \text{and} \quad \frac{\frac{\vdash \Gamma_2}{\vdash \Gamma, \perp} \perp \quad \frac{\vdash \Gamma_1}{\vdash \Gamma, \perp} \perp}{\vdash \Gamma, \perp \& \perp} \&$$

are interpreted by variants, i.e. are not distinguished. In fact proposition 5 would become false. In terms of  $\flat$ -calculus, a new problematic case arises, namely a cut

$$\frac{\frac{}{\vdash 1, \flat} \flat \quad \frac{\vdash \Gamma}{\vdash \perp, \Gamma} \perp}{\vdash \Gamma, \flat} CUT$$

for which the only natural cut-elimination would be to add a rule of weakening on  $\flat$  (from  $\vdash \Gamma$ , deduce  $\vdash \Gamma, \flat$ ) which contradicts our idea that a flat proof is a proof in which something is actually missing : this rule would acknowledge ordinary proofs as flat ones, thus destroying all our construction. In other terms, the “natural” geometry of interaction for  $\perp$  leads to accept the rule : from  $\vdash \Gamma$  and  $\vdash \Gamma$  deduce  $\vdash \Gamma$ , which is a strong form of summation. With such a rule, booleans would no longer be standard (any formal sum of booleans would be accepted).

Instead, we modify the  $\perp$ -rule into : from  $\vdash A, \Gamma$  deduce  $\vdash \perp, A, \Gamma$ , in which one of the formulas in the premise is distinguished. Strictly speaking, there is no need to distinguish  $A$ , if we stick to the viewpoint of sequents as sequences (instead of multisets). With this modified rule, the problem of a cut

$$\frac{\frac{\text{---} \flat \quad \frac{\vdash A, \Gamma}{\vdash \perp, A, \Gamma} \perp}{\vdash 1, \flat} \text{---} \text{CUT}}{\vdash A, \Gamma, \flat}$$

disappears ; for instance the cut can be replaced with a cut between  $\vdash A, \Gamma$  and the flat axiom  $\vdash A^\perp, A, \flat \dots$

### Definition 16

*The rules for the multiplicative units are interpreted as follows :*

- *If  $\Pi$  is the axiom  $\vdash 1$ , then  $\Pi^\bullet$  is the message  $p_1(x, y)$  ;*
- *If the proof  $\Pi$  of  $\vdash [\Delta], \perp, A, \Gamma$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], A, \Gamma$  by a  $\perp$ -rule, then  $\Pi^\bullet$  is defined in terms of  $\Pi_1^\bullet$  :*
  - *in  $\Pi_1^\bullet$  replace any wire  $p_B(t, u) \mapsto p_A(t', u)$  with a wire  $p_B(t, u) \mapsto p_\perp(t', u)$  ;*
  - *add to the result of this replacement the wire  $p_\perp(x, y) \mapsto p_A(x, y)$ .*

### Remark 14

- One should prove the analogue of theorem 4 which can be done without difficulty. In case of extreme laziness, observe that a (slightly more complicated) geometry of interaction of multiplicative neutrals can easily be given by means of the second-order translation of  $\perp$  as  $\exists \alpha(\alpha \otimes \alpha^\perp)$ , for which the previous section yields a geometry of interaction. Our modified  $\perp$ -rule can be translated in second-order linear logic by

$$\frac{\frac{\text{---} ID}{\vdash \Gamma, A \quad \vdash A^\perp, A} \otimes}{\vdash \Gamma, A \otimes A^\perp, A} \exists$$

$$\vdash \Gamma, \exists \alpha(\alpha \otimes \alpha^\perp), A$$

- In terms of proof-nets, the “natural choice” for  $\perp$  corresponds to proof-nets where the weakened formula ( $\perp$  or  $?A$ ) is physically disconnected. There is no reasonable correctness criterion for such nets :
  - In the case of a multiplicative formula  $A$  using only multiplicative units as atoms, such a choice would lead to identify all proofs of  $A$  ;
  - Hence the correctness problem for such proof-nets contains the decision problem for the constant-only multiplicative fragment, which is known to be **NP**-complete, see [LW92].
  - But the general shape of the known criterions is **coNP**<sup>8</sup>, hence the existence of a criterion of the same shape is very unlikely...

Our solution corresponds to a version of proof-nets in which weakened formulas are attached to a formula of the net.

- The geometry of interaction of  $\perp$  (which is the geometry of the weakening rule) shows for the first time a non-contrived non-commutativity : the formula to the left of  $\perp$  in the conclusion of the rule actually matters. One could imagine another version in which  $\perp$  is physically linked to its rightmost neighbor. This non-commutative reading is controversial anyway : we could also decide to say that the rule involves two formulas,  $\perp$  and  $A$ .
- The neutrals display another originality w.r.t. geometry of interaction : in the absence of  $\perp$  and  $1$ ,  $\Pi^*$  is always a sum of two adjoint wirings  $W + W^*$ . The interpretation of the axiom for  $1$  is a single wire, and the  $\perp$ -rule introduces non self-adjoint wirings...

The additive neutrals will be easily interpreted, provided the usual axiom for  $\top$  is replaced with the rule

$$\frac{\vdash \Gamma, \flat}{\vdash \Gamma, \top} \top$$

When  $\Pi$  is obtained from  $\Sigma$  by means of a  $\top$ -rule, then

$$\Pi^\bullet = \Sigma^\bullet + p_\top(x, y)$$

This definitely clarifies the relation between  $\top$  and  $\flat$  ; the addition of the message  $p_\top(x, y)$  ensures  $\Pi^\bullet \neq 0$ . The cut-elimination procedure is extended in the following way : a cut between the flat axiom  $\vdash \Phi, 0, \flat$  and the conclusion  $\vdash \Psi, \top$  of a  $\top$ -rule is replaced with the proof of  $\vdash \Psi, \Phi, \flat$  obtained from the premise  $\vdash \Psi, \flat$  of the  $\top$ -rule by “inflation” (see remark 12). Soundness is almost immediate.

---

<sup>8</sup>Typically :  $\pi$  is a multiplicative proof-net iff for all switchings  $\mathcal{S}$  the resulting graph is connected and acyclic ; this **coNP** turns out to be polytime, more precisely quadratic.

## 4 The case of exponentials

### 4.1 An alternative version of exponentials

We describe here a modification of the exponential rules ; the modification is close to some extant experimental systems, see e.g. [A93, MM93]. The system has a weak form of promotion, which is corrected by an additional rule for  $?$ .

$$\frac{\vdash \Gamma, A}{\vdash ?\Gamma, !A} ! \quad \frac{\vdash A, \Gamma}{\vdash ?A, \Gamma} d? \quad \frac{\vdash B, \Gamma}{\vdash ?A, B, \Gamma} w? \quad \frac{\vdash ?A, ?A, \Gamma}{\vdash ?A, \Gamma} c? \quad \frac{\vdash ??A, \Gamma}{\vdash ?A, \Gamma} ??$$

These rules are respectively called (*weak*) *promotion*, *dereliction*, *contraction*, *digging*. They only ensure a variant of the subformula property ( $??A$  subformula of  $?A$ ) which may look strange at first sight, but which is not more artificial than the familiar definition which allows  $A[t/x]$  to be a subformula of  $\forall x A$ . This modification of the notion of subformula yields infinitely many subformulas for a propositional linear formula, in accordance with the undecidability of propositional linear logic, [LMSS90]. The rule for weakening explicitly mentions a formula  $B$  of the context, in accordance with what we did for  $\perp$ . The rule of promotion is understood in the obvious way : if  $\Gamma$  is  $A_1, \dots, A_n$ , then  $?\Gamma$  is  $?A_1, \dots, ?A_n$ . We are therefore interpreting a slight modification of linear logic, corresponding to extant experimental systems. If we adopt here the modifications of [A93, MM93], this is because geometry of interaction is particularly simple in this setting. But also, since geometry of interaction is in many senses the most natural semantics, it might be seen as backing the syntactical variants proposed in these works.

**Exercise 2** *Prove the cut-elimination theorem for this variant of the exponential rules. Of course one has to define an ad hoc cut-elimination procedure.*

### 4.2 The pattern

The general pattern is as follows : we want to allow reuse ; the absolutely weakest form of reuse is expressed by the principle  $\vdash ?A^\perp, A \otimes A$ . The interpretation  $\Pi^\bullet$  of the (natural) proof of this principle must have the following property : if  $U$  interprets a proof of  $A$  and  $!U$  is the result of the promotion of  $U$ , then a cut between  $!U$  and  $\Pi^\bullet$  yields after normalisation a variant of  $V$ , where  $V$  is obtained by means of definition 10.4 from  $U$  and  $U$  in the respective roles of  $\Pi_1^\bullet$  and  $\Pi_2^\bullet$ . In the formation of  $V$  the essential step is the formation of  $\otimes_1(U) + \otimes_2(U)$ , which is the sum of two variants. The execution formula is able to extract two variants of  $U$  from the sole  $!U$ . Building variants basically involves non-laminated wirings, and therefore the execution formula is unable to achieve the task. In particular the solution  $!U = U$  is inadequate.

Of course, the problem is easily solved if we could drop lamination. . . , but this would destroy some essential features of our construction. But we can also mimic non-laminated wires by laminated ones : we are eventually lead to shift the dialectal components  $ux$  of  $\otimes_1(U)$  from the private part to the public part. The context-free promotion  $!U$  is interpreted as follows : in  $U$  replace all atoms  $p_A(t, u)$  with atoms  $p_{!A}(txu, z)$  ; This construction is the result of two steps :

- We first form  $\otimes_1(U)$ , by replacing in  $U$  all atoms  $p_B(t, u)$  with atoms  $p_B(t, xu)$ .
- Then we replace in  $\otimes_1(U)$  all atoms  $p_B(t, xu)$  with atoms  $p_{!A}(txu, z)$ . The final  $z$  is present only because our predicates need a second argument : in reality there is no dialect in this case.

### 4.3 Interpretation of the exponential rules

In what follows  $t_1 \dots \hat{t}_i \dots t_n$  is short for  $t_1 \dots t_{i-1} t_{i+1} \dots t_n$ .

**Definition 17 (def. 10 cont<sup>d</sup>)**

**Case 17.1** Assume that the proof  $\Pi$  of  $\vdash [\Delta], ?\Gamma, !A$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], \Gamma, A$  by a promotion rule ; we define  $\Pi^\bullet$  in terms of  $\Pi_1^\bullet$ . Let us assume that  $\Gamma$  is  $C_1, \dots, C_n$  (in this order). Replace in  $\Pi_1^\bullet$  :

- all atoms  $p_A(t, u)$  with atoms  $p_{!A}(txuy_1 \dots y_n, z)$  ;
- all atoms  $p_B(t, u)$  with atoms  $p_B(txuy_1 \dots y_n, z)$  when  $B$  is in  $\Delta$  ;
- all atoms  $p_{C_i}(t, u)$  with atoms  $p_{?C_i}(t(uy_1 \dots \hat{y}_i \dots y_n x)y_i, z)$ .

( $x, y_1, \dots, y_n, z$  are fresh variables). as usual  $x, x', x'', z$  are fresh variables. The result is by definition  $\Pi^\bullet$ .

**Case 17.2** Assume that the proof  $\Pi$  of  $\vdash [\Delta], ?A, \Gamma$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], A, \Gamma$  by a dereliction rule ; we define  $\Pi^\bullet$  by replacing in  $\Pi_1^\bullet$  :

- All wires  $p_C(t, u)$  with atoms  $p_{?A}(tgz, uz)$  ;
- All atoms  $p_B(t, u)$  (when  $B \neq A$ ) with atoms  $p_B(t, uz)$ .

**Case 17.3** Assume that the proof  $\Pi$  of  $\vdash [\Delta], ?A, B, \Gamma$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], B, \Gamma$  by a weakening rule ; we define  $\Pi^\bullet$  in terms of  $\Pi_1^\bullet$ .

- in  $\Pi_1^\bullet$  replace any wire  $p_C(t, u) \mapsto p_B(t', u)$  with a wire  $p_C(t, u(xx'y)) \mapsto p_{?A}(xt'y, u(xx'y))$  ;

- in  $\Pi_1^\bullet$  replace any wire  $p_C(t, u) \mapsto p_D(t', u)$ , with  $D \neq B$  with a wire  $p_C(t, u(xx')y) \mapsto p_D(t', u(xx')y)$  ;
- add to the result of this replacement the wire  $p_{?A}(x'zy', z'(xx')y) \mapsto p_B(z, z'(xx')y)$ .

**Case 17.4** Assume that the proof  $\Pi$  of  $\vdash [\Delta], ?A, \Gamma$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], ?A_1, ?A_2, \Gamma$  by a contraction rule (we note  $?A_1, ?A_2$  two distinct occurrences of  $A$ ) ; we define  $\Pi^\bullet$  in terms of  $\Pi_1^\bullet$ .

- First put the atoms  $p_{?A_i}(t, u)$  “in the form  $p(tt't'', u)$ ” ; this means that we form

$$W = \sum_{i=1}^2 p_{?A_i}(xx'x'', y) + \sum_{C \in \Gamma} p_C(x, y)$$

and replace  $\Pi_1^\bullet$  with  $U = W\Pi_1^\bullet W$ .

- Then we replace in  $U$  :
  - All atoms  $p_{?A_1}(tt't'', u)$  with atoms  $p_{?A}(t(gt')t'', u)$  ;
  - All atoms  $p_{?A_2}(tt't'', u)$  with atoms  $p_{?A}(t(dt')t'', u)$ .

the result is by definition  $\Pi^\bullet$ .

**Case 17.5** Assume that the proof  $\Pi$  of  $\vdash [\Delta], ?A, \Gamma$  is obtained from a proof  $\Pi_1$  of  $\vdash [\Delta], ??A, \Gamma$  by a digging rule ; we define  $\Pi^\bullet$  in terms of  $\Pi_1^\bullet$ .

- First put the atoms  $p_{??A}(t, u)$  “in the form  $p((tt't'')uu', v)$ ” ; this means that we form

$$W = p_{??A}((xx'x'')yy', z) + \sum_{C \in \Gamma} p_C(x, y)$$

and replace  $\Pi_1^\bullet$  with  $U = W\Pi_1^\bullet W$ .

- Then we replace in  $U$  all atoms  $p_{??A}((tt't'')uu', v)$  with atoms  $p_{?A}(t(t'uu')t'', v)$  ;

## Remark 15

The rule of promotion is the first violation of the principle that all our constructions are compatible with  $\sim$ . This is perhaps the ultimate meaning of the “!-box”.

## 4.4 Basic examples

We shall treat certain basic cuts

$$\frac{\frac{\dots}{\vdash \Gamma, !A} R \quad \frac{\dots}{\vdash ?A^\perp, \Delta} S}{\vdash \Gamma, \Delta} CUT$$

on an exponential  $!A$ , in certain cases :

- $\Gamma$  is empty and  $R$  is the promotion rule applied to a cut-free proof  $\Pi$  of  $\vdash A$ .
- $S$  is a  $!, d?, w?, c?$  or  $??$ -rule, applied to a cut-free proof  $\Sigma$  of  $\vdash \Delta'$  and with main formula  $?A^\perp$  in case  $S \neq !$ .

### Example 2

Cut-elimination replaces the cut

$$\frac{\frac{\vdash A}{\vdash !A} ! \quad \frac{\vdash A^\perp, B}{\vdash ?A^\perp, !B} !}{\vdash !B} CUT$$

with

$$\frac{\frac{\vdash A \quad \vdash A^\perp, B}{\vdash B} CUT}{\vdash !B} !$$

The interpretation of the original proof is  $(U + V, \sigma_{!A, ?A^\perp; !B})$ , where :

- $U$  is obtained by replacing in  $\Pi^\bullet$  all atoms  $p_A(t, u)$  with atoms  $p_{!A}(txu, zz')$  ;
- $V$  is obtained by replacing in  $\Sigma^\bullet$  :
  - all atoms  $p_{A^\perp}(t, u)$  with atoms  $p_{?A^\perp}(t(ux)y, zz')$
  - all atoms  $p_B(t, u)$  with atoms  $p_{!B}(txuy, zz')$ .

The interpretation of the modified proof is  $(W + Y, \sigma_{A, A^\perp; !B})$ , where :

- $W$  is obtained by replacing in  $\Pi^\bullet$  all atoms  $p_A(t, u)$  with atoms  $p_A(txyu, z)$  ;
- $Y$  is obtained by replacing in  $\Sigma^\bullet$  all atoms  $p_{A^\perp}(t, u)$  with atoms  $p_{A^\perp}(txuy', z)$  and all atoms  $p_B(t, u)$  with atoms  $p_{!B}(txuy', z)$ .



Observe that :

- $RES(U + V, \sigma_{!A, ?A^\perp; !B}) = RES(U_1 + V, \sigma_{!A, ?A^\perp; !B})$ , where  $U_1$  is obtained from  $U$  by replacing all atoms  $p_{!A}(txu, zz')$  with atoms  $p_{!A}(t(x'x'')u, zz')$  ;
- $RES(U_1 + V, \sigma_{!A, ?A^\perp; !B}) = RES(Z, \sigma_{A, A^\perp; !B})$ , where  $Z$  is obtained from  $U_1 + V$  by replacing :
  - all atoms  $p_{!A}(t(x'x'')u, zz')$  with atoms  $p_A(tx''x'u, zz')$
  - all atoms  $p_{?A^\perp}(t(ux)y, zz')$  with atoms  $p_{?A^\perp}(tuxy, zz')$ .
- $Z = \otimes_1(W + Y)$ .

therefore  $RES(U + V, \sigma_{!A, ?A^\perp; !B}) = \otimes_1(RES(W + Y, \sigma_{A, A^\perp; !B}))$ . This shows the soundness of this particular cut-elimination step.

### Exercise 3

Extend example 2 to the more general case of  $n$  cuts between :

- Cut-free proofs of the sequents  $\vdash !A_1, \dots, \vdash !A_n$
- A cut free proof of the sequent  $\vdash !A_1, \dots, !A_n$  ending with a promotion rule.

### Example 3

Cut-elimination replaces the cut

$$\frac{\frac{\vdash A}{\vdash !A}! \quad \frac{\vdash A^\perp, \Phi}{\vdash ?A^\perp, \Phi}d?}{\vdash \Phi}CUT$$

with

$$\frac{\vdash A \quad \vdash A^\perp, \Phi}{\vdash \Phi}CUT$$

The interpretation of the original proof is  $(U + V, \sigma_{!A, ?A^\perp; \Phi})$ , where  $U$  is as in example 2 ; and  $V$  is obtained by replacing in  $\Sigma^\bullet$

- All atoms  $p_A^\perp(t, u)$  with atoms  $p_{?A^\perp}(tgz, (uz)z')$  ;
- All atoms  $p_B(t, u)$  (when  $B \neq A^\perp$ ) with atoms  $p_B(t, (uz)z')$ .

whereas the modified proof is interpreted by  $(\otimes_1(\Pi^\bullet) + \otimes_2(\Sigma^\bullet), \sigma_{A, A^\perp; \Phi})$ . Now observe that :

- $RES(U + V, \sigma_{!A, ?A^\perp; \Phi}) = RES(U + V, \sigma_{!A, ?A^\perp; \Phi})$
- $RES(U + V, \sigma_{!A, ?A^\perp; \Phi}) = RES(U_1 + V, \sigma_{!A, ?A^\perp; \Phi})$ , where  $U_1$  is obtained from  $U$  by replacing all atoms  $p_{!A}(txu, zz')$  with atoms  $p_{!A}(tgu, (z_1u)z')$  ;
- $RES(U_1 + V, \sigma_{!A, ?A^\perp; \Phi}) = RES(W, \sigma_{A, A^\perp; \Phi})$ , where  $W$  is obtained from  $U_1 + V$  by replacing all atoms  $p_{!A}(tgu, (z_1u)z')$  with atoms  $p_A(t, (z_1u)z')$  and all atoms  $p_{?A^\perp}(tgz, (uz)z')$  with atoms  $p_{A^\perp}(t, (uz)z')$  ;
- $W = \otimes_2(\otimes_1(\Pi^\bullet) + \otimes_2(\Sigma^\bullet))$ .

therefore  $RES(U + V, \sigma_{!A, ?A^\perp; \Phi}) = \otimes_2(RES(\otimes_1(\Pi^\bullet) + \otimes_2(\Sigma^\bullet), \sigma_{A, A^\perp; \Phi}))$ . This shows the soundness of this particular cut-elimination step.

#### Example 4

Cut-elimination replaces the cut

$$\frac{\frac{\frac{}{\vdash A}! \quad \frac{}{\vdash B, \Phi} w?}{\vdash !A \quad \vdash ?A^\perp, B, \Phi}}{\vdash B, \Phi} CUT$$

with the proof  $\Sigma$ . The interpretation of the original proof is  $(U + V, \sigma_{!A, ?A^\perp; \Phi})$ , where  $U$  is as in example 2 ; and  $V$  is defined in terms of  $\Sigma^\bullet$  :

- In  $\Sigma^\bullet$  replace any wire  $p_C(t, u) \mapsto p_B(t', u)$  by a wire  $p_C(t, (u(xx')y)z'') \mapsto p_{?A^\perp}(xt'y, (u(xx')y)z'')$  ;
- in  $\Sigma^\bullet$  replace any wire  $p_C(t, u) \mapsto p_D(t', u)$ , with  $D \neq B$  by a wire  $p_C(t, (u(xx')y)z'') \mapsto p_D(t', (u(xx')y)z'')$  ;
- add to the result of this replacement the wire  $p_{?A^\perp}(x'zy', (z'(xx')y)z'') \mapsto p_B(z, (z'(xx')y)z'')$ .

Now observe that nilpotency  $\sigma_{!A, ?A^\perp; \Phi}(U + V)$  is immediate, and that  $RES(U + V, \sigma_{!A, ?A^\perp; \Phi})$  is the sum of the following wires :

- all wires  $p_C(t, (u(xx')y)z'') \mapsto p_D(t', (u(xx')y)z'')$ , where  $p_C(t, u) \mapsto p_D(t', u)$  is a wire in  $\Sigma^\bullet$  and  $D \neq B$
- all wires  $p_C(t, (u(ww')v)z'') \mapsto p_B(t', (u(ww')v)z'')$ , where  $p_C(t, u) \mapsto p_B(t', u)$  is a wire in  $\Sigma^\bullet$ , and  $p_{!A}(wxv, zz') \mapsto p_{!A}(w'xv, zz')$  is a wire in  $U$ .

This wiring is easily shown to be a variant of  $\Sigma^\bullet$  ; however, the fact that  $\Sigma^\bullet \neq 0$  plays an essential role : without at least one wire  $p_{!A}(wxv, zz') \mapsto p_{!A}(w'xv, zz')$  in  $U$ ,

$RES(U + V, \sigma_{!A, ?A^\perp, \Phi})$  cannot keep any track of the wires  $p_C(t, u) \mapsto p_B(t', u)$  of  $\Sigma^\bullet$ . This shows the soundness of this particular cut-elimination step.

### Example 5

Cut-elimination replaces the cut

$$\frac{\frac{\frac{}{\vdash A}!}{\vdash !A} \quad \frac{\vdash ?A_1^\perp, ?A_2^\perp, \Phi}{\vdash ?A^\perp, \Phi} c?}{\vdash \Phi} CUT$$

with

$$\frac{\frac{\frac{}{\vdash A}!}{\vdash !A} \quad \frac{\frac{\frac{}{\vdash A}!}{\vdash !A_1} \quad \vdash ?A_1^\perp, ?A_2^\perp, \Phi}{\vdash ?A_2^\perp, \Phi} CUT}{\vdash \Phi} CUT$$

The interpretation of the original proof is  $(U + V, \sigma_{!A, ?A^\perp, \Phi})$ , where  $U$  is as in example 2 ; and  $V$  is obtained from  $\Sigma^\bullet$  as follows :

- First define  $V' = \otimes_2(\Sigma^\bullet)$  ;
- First put the atoms  $p_{?A_i}(t, u)$  of  $V'$  “in the form  $p(tt't'', u)$ ” ; this means that we form as in definition 17.4

$$W = p_{?A_1}(xx'x'', y) + p_{?A_2}(xx'x'', y) + \sum_{C \in \Phi} p_C(x, y)$$

and replace  $V'$  with  $V'' = WV'W$  ;

- Then we replace in  $V''$  :
  - All atoms  $p_{?A_1}(tt't'', u)$  with atoms  $p_{?A}(t(gt')t'', u)$  ;
  - All atoms  $p_{?A_2}(tt't'', u)$  with atoms  $p_{?A}(t(dt')t'', u)$ .

the result is called  $V$ .

The interpretation of the modified proof is  $(X_1 + X_2 + Y, \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp, \Phi})$  where :

- $X_1 = U_1$

- $X_2 = \otimes_2(U_2)$
- $Y = \otimes_2(\otimes_2(\Sigma^\bullet))$

where  $U_i$  is obtained from  $U$  by replacing the predicate  $p_{!A}$  by  $p_{!A_i}$ . Now observe that  $U$  (which is of the form  $\otimes_1(U')$  is also of the form  $\otimes_2(\otimes_2(U''))$ .  $X_1 + X_2 + Y$  is therefore of the form  $\otimes_2(\otimes_2(Z))$ , which is a variant of  $\otimes_2(Z)$  by proposition 2. Therefore by proposition 1 we can take as interpretation of our modified proof the pair  $(X'_1 + X'_2 + Y', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi})$ , where :

- $X_1 = \otimes_2(X'_1)$
- $X'_2 = U_2$
- $Y = \otimes_2(\Sigma^\bullet)$

But it is immediate to see that

$$RES(X'_1 + X'_2 + Y', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi}) = RES(X''_1 + X'_2 + Y', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi})$$

where  $X''_1$  is obtained from  $X'_1$  by replacing all atoms  $p_{A_1}(t, z)$  with atoms  $p_{A_1}(t, zz')$  ; but observe that  $X''_1 = \otimes_2(X'_1)$ , hence  $X''_1 = X_1 = U_1$ . Now observe that all the wires in  $X_1$  and  $X'_2$  are of the form  $p_{A_i}(tt't'', u)$ , hence

$$RES(X_1 + X'_2 + Y', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi}) = RES(X_1 + X'_2 + Y'', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi}),$$

where  $Y''$  is obtained by putting the atoms  $p_{?A_i}(t, u)$  of  $Y'$  “in the form  $p(tt't'', u)$ ” ; but then  $Y'' = V''$ . We are left with

$RES(U_1 + U_2 + V'', \sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi})$ . Obviously this expression is unchanged if we merge the  $p_{?A_i}$  into  $p_{?A}$  and the  $p_{!A_i}$  into  $p_{!A}$ . This merge changes  $V''$  into  $V$ ,  $U_1 + U_2$  into  $ZUZ$  and  $\sigma_{!A_1, ?A_1^\perp, !A_2, ?A_2^\perp; \Phi}$  into  $Z\sigma_{!A, ?A^\perp; \Phi}Z$ , where  $Z$  is the projection

$$Z = p_{?A}(x(gx')x'', y) + p_{?A}(x(dx')x'', y) + \sum_{C \in \Gamma} p_C(x, y)$$

Observe that  $V = ZVZ$ , hence

$$\begin{aligned} RES(ZUZ + V, Z\sigma_{!A, ?A^\perp; \Phi}Z) &= Z \cdot RES(U + V, \sigma_{!A, ?A^\perp; \Phi}) \cdot Z = \\ &= RES(U + V, \sigma_{!A, ?A^\perp; \Phi}) \end{aligned}$$

and thus we have eventually proved soundness.

### Example 6

Cut-elimination replaces the cut

$$\frac{\frac{\vdash A}{\vdash !A}! \quad \frac{\vdash ??A^\perp, \Phi}{\vdash ?A^\perp, \Phi}??}{\vdash \Phi} CUT$$

with

$$\frac{\frac{\frac{\vdash A}{!}}{\vdash !A}{!}}{\vdash !!A} \quad \vdash ??A^\perp, \Phi \quad \frac{}{CUT} \quad \frac{}{\vdash \Phi}$$

The interpretation of the original proof is  $(U + V, \sigma_{!A, ??A^\perp, \Phi})$ , where  $U$  is as in example 2 ; and  $V$  is obtained from  $\Sigma^\bullet$  as follows :

- First put the atoms  $p_{??A}(t, u)$  “in the form  $p((tt't'')uu', v)$ ” ; this means that we form

$$W = p_{??A}((xx'x'')yy', z) + \sum_{C \in \Phi} p_C(x, y)$$

as in definition 17.5 and replace  $\Pi_1^\bullet$  with  $V' = W\Sigma^\bullet W$ .

- Then we replace in  $V'$  all atoms  $p_{??A}((tt't'')uu', v)$  with atoms  $p_{?A}(t(t'uu')t'', vy)$  ;

the result is by definition V. The modified proof is interpreted as  $(X + Y, \sigma_{!!A, ??A^\perp, \Phi})$ , with :

- $X$  is obtained from  $\Pi^\bullet$  by replacing all atoms  $p_A(t, u)$  with atoms  $p_{!!A}((txu)yz, z')$  ;
- $Y = \otimes_2(\Sigma^\bullet)$

We first observe that  $RES(X + Y, \sigma_{!!A, ??A^\perp, \Phi}) = RES(X' + \otimes_2(V'), \sigma_{!A, ?A^\perp, \Phi})$ , where  $X'$  is obtained from  $\Pi^\bullet$  by replacing all atoms  $p_A(t, u)$  with atoms  $p_{!A}(t(xyz)u, z')$ . This is because  $X' = ZXZ^*$  and  $\otimes_2(V') = ZYZ^*$  with

$$Z = p_{!!A}((xx'x'')yy', z) \mapsto p_{!A}(x(x'yy')x'', z) + p_{??A^\perp}((xx'x'')yy', z) \mapsto p_{??A^\perp}(x(x'yy')x'', z) + \sum_{C \in \Phi} p_C(x, y)$$

Now observe that

$$RES(X' + \otimes_2(V'), \sigma_{!A, ?A^\perp, \Phi}) = RES(X' + WVW, \sigma_{!A, ?A^\perp, \Phi})$$

but

$$\begin{aligned} RES(X' + WVW, \sigma_{!A, ?A^\perp, \Phi}) &= RES(WX'W + V, \sigma_{!A, ?A^\perp, \Phi}) = \\ &= RES(U + V, \sigma_{!A, ?A^\perp, \Phi}) \end{aligned}$$

This proves soundness in this last case.

## 4.5 Nilpotency for exponentials

It remains to extend the nilpotency theorem to the full case ; this offers no difficulty of principle (basically our previous treatment of exponentials in [G88] suitably modified in the spirit of section 3). For instance the nilpotency theorem basically needs an extension of definition 12 :

**Definition 18 (def. 12 cont<sup>d</sup>)**

**Case 18.1** *If  $X$  is a weak type, then we define the weak type  $!X$  as  $Z^{\perp\perp}$ , where  $Z$  consists of all  $!U$  obtained from some  $U \in X$  by means of definition 17.1 (in the simplified case  $\Gamma = \Delta = \emptyset$ ).*

**Case 18.2** *If  $X$  is a weak type, then we define the weak type  $?X$  as  $(!X^\perp)^\perp$ .*

This definition is the key to an unproblematic extension of the nilpotency theorem to the full case (left without hypocrisy to the reader).

**Theorem 5 (Nilpotency)**

$\sigma\Pi^\bullet$  is nilpotent.

## 4.6 Soundness for exponentials, a sketch

Soundness is more delicate ; in fact we can only prove soundness under a strong restriction on  $\Gamma$ .

**Theorem 6 (Limited soundness)**

*If a cut-free proof  $\Pi$  of  $\vdash \Gamma$  is obtained from a proof  $\Sigma$  of  $\vdash [\Delta]\Gamma$  by means of the transformations sketched in definition 15 (suitably extended to accommodate exponentials, see below), and if  $\Gamma$  contains no exponential and no second order existential quantifier, then  $\Pi^\bullet$  is a variant of  $RES(\Sigma^\bullet, \sigma_{\Delta;\Gamma})$ .*

PROOF: the proof is an imitation of theorem 1 ii) of [G88]. We give some hints :

- The restriction on  $\Gamma$  makes it possible to consider a limited form of cut-elimination, where an exponential cut is eliminated only when the premise containing  $!A$ , let us say  $\Phi, !A$ , comes from a  $!$ -rule, with  $\Phi$  empty. One can show that cut-elimination works with this restricted algorithm.
- The examples 2 (including its n-ary generalization of exercise 3), 3, 4, 5, 6 are the basic steps of the verification of the soundness of this limited cut-elimination algorithm.

- Among the specific technicalities of this extension, we need to develop the  $\flat$ -calculus in presence of exponentials. For instance the promotion rule in presence of  $\flat$  is understood as follows : since  $?b$  is an illegal expression we can form  $? \Gamma$  only when  $\Gamma$  does not contain  $\flat$ . One also needs to define cut-elimination between a flat axiom and an exponential rule. This is straightforward :
  - A cut between the flat axiom  $\vdash \Gamma, !A, \flat$  and a  $d?$ -rule is replaced with a cut between the flat axiom  $\vdash \Gamma, A, \flat$  and the premise of the  $d?$ -rule ;
  - A cut between the flat axiom  $\vdash \Gamma, !A, \flat$  and a  $w?$ -rule, with premise  $\vdash B, \Gamma$  is replaced with a cut between the flat axiom  $\vdash \Gamma, B^\perp, \flat$  and the premise of the  $w?$ -rule ;
  - A cut between the flat axiom  $\vdash \Gamma, !A, \flat$  and a  $c?$ -rule is replaced with two cuts between the flat axioms  $\vdash \Gamma, !A, \flat$  and  $\vdash !A, \flat$  and the premise of the  $c?$ -rule ;
  - A cut between the flat axiom  $\vdash \Gamma, !A, \flat$  and a  $??$ -rule is replaced with a cut between the flat axiom  $\vdash \Gamma, !!A, \flat$  and the premise of the  $??$ -rule ;
  - A cut between a promotion  $!A$  and the flat axiom  $\vdash \Gamma, ?A^\perp, \flat$  is replaced with the flat axiom  $\vdash \Gamma, b$ .  $\square$

### Remark 16

- One should also prove the soundness w.r.t. the full cut-elimination procedure. The only way to do so would be to work with *proof-nets* in some variant of [G94] (with boxes for  $!$ ), then to prove a Church-Rosser property, and the fact that the interpretation of a cut-free proof of  $\Gamma$  depends (up to variance) only on the associated net. This seems to be unproblematic, but we quailed in the face of the burden. . .
- In [G88] we were able to prove a little more, namely  $!$  was allowed in  $\Gamma$ . The problem here is that reduction above a promotion rule would need a more general notion of variant. If we come back to definition 6, we see that the essential point is the possible choices for  $W$  and  $W'$ . We can perfectly define, *for each type  $A$*  a set of wirings (inducing an equivalence  $\sim_A$ , coarser than  $\sim$ ). For instance for  $!A$  we could consider all wirings  $\sum p_{!A}(xyt, z) \mapsto p_{!A}(xyu, z)$  to enlarge the notion of variant in that case. No doubt that with this extension, we can allow  $!$  in  $\Gamma$ .
- What about a complete soundness ? Surely, one needs a liberalized notion of variant, as above, typically to take care of rather arbitrary choices

made in the case of the  $\text{?}$ -rules, but is this enough ? Presumably not, and we honestly don't know whether the  $\text{b}$ -calculus (perhaps improved with additional principles) can cope with the situation. In fact we have conflicting intuitions :

- Dr Jekyll thinks that the syntax can be adapted to cope with the geometrical interpretation ;
  - Mr Hyde thinks that there is something basically infinite in exponentials and that for this reason, there is an irreducible global configuration, the  $\text{!}$ -box.
- Anyway, we do not want to achieve soundness by replacing the rather natural notion of being variants (with the possibility of modifying the possible  $W, W'$ ) with a notion of observational equivalence : the definition of variance should remain rather elementary, if possible.

## 5 Appendix : Hilbert spaces and related topics

### 5.1 Hilbert spaces

#### Definition 19

A prehilbertian space is a complex vector space  $\mathcal{H}$  equipped with a positive hermitian form, i.e. a function  $(x \mid y)$  from  $\mathcal{H} \times \mathcal{H}$  to  $\mathbb{C}$  such that :

- $(\alpha x + \beta y \mid z) = \alpha(x \mid z) + \beta(y \mid z)$ , for  $x, y, z \in \mathcal{H}$ ,  $\alpha, \beta \in \mathbb{C}$
- $(x \mid y) = \overline{(y \mid x)}$  for  $x, y \in \mathcal{H}$  ; in particular  
 $(z \mid \alpha x + \beta y) = \bar{\alpha}(z \mid x) + \bar{\beta}(z \mid y)$  and  $(x \mid x)$  is always real.
- $(x \mid x) \geq 0$  for  $x \in \mathcal{H}$

Among the immediate properties of such spaces, let us mention the famous Cauchy-Schwarz inequality :  $|(x \mid y)|^2 \leq (x \mid x)(y \mid y)$ , which implies that  $\|x\| = (x \mid x)^{1/2}$  defines a semi-norm on  $\mathcal{H}$ . Another classic is the “median identity”  $\|x + y\|^2 + \|x - y\|^2 = 2\|x\|^2 + 2\|y\|^2$ .

#### Definition 20

$\mathcal{H}$  is said to be a Hilbert space when  $\|\cdot\|$  is a norm, i.e. when  $(x \mid x) > 0$  for all  $x \neq 0$  and  $\mathcal{H}$  is complete (i.e. is a so-called Banach space) w.r.t.  $\|\cdot\|$ .

Every prehilbertian space  $\mathcal{H}$  can be transformed into a Hilbertian space  $\mathcal{H}''$  ; the process involves two steps :



- Separation : quotient  $\mathcal{H}$  by the subspace consisting of vectors with a null semi-norm. This induces a hermitian form on the quotient  $\mathcal{H}'$ , which is a norm, i.e. is such that  $\|x\| = 0$  implies  $x = 0$ .
- Completion : add limits for all Cauchy sequences, and extend the hermitian form to this extended space  $\mathcal{H}''$ .

### Example 7

Let  $I$  be a set ; we define  $\ell^2(I)$  to consist of all square-summable sequences of complex numbers indexed by  $I$ , i.e. of all families  $(\lambda_i)_{i \in I}$  (sometimes noted  $\Sigma \lambda_i.i$ ) such that  $\sum_{i \in I} |\lambda_i|^2 < +\infty$ . We define a hilbertian form on  $\mathcal{H}$  by  $((\lambda_i) | (\mu_i)) = \sum_{i \in I} \lambda_i \overline{\mu_i}$  (the series is shown to be absolutely convergent by a direct proof of the Cauchy-Schwarz inequality).  $\ell^2(I)$  is easily shown to be a Hilbert space ; in fact general (and quite easy) results on Hilbert space shows that any Hilbert space is isomorphic to some  $\ell^2(I)$ . Since the isomorphism class of  $\ell^2(I)$  only depends on the cardinality of  $I$ , we see that there are three cases :

- If  $I$  is finite, then the vector space is finite dimensional ; such Hilbert spaces are too small in practice
- If  $I$  is infinite but not denumerable, then the space is too big for most applications.
- If  $I$  is denumerable, then we get the main Hilbert space, “the” Hilbert space. It must be noticed that although in practice most Hilbert spaces will fall in this equivalence class, the isomorphism might be non-trivial, i.e. there is basically one space, but it may appear through very unlikely disguises.

### Definition 21

*If  $\mathcal{H}$  and  $\mathcal{H}'$  are Hilbert spaces, then we can form a new space  $\mathcal{H}'' = \mathcal{H} \oplus \mathcal{H}'$  by considering the set of all formal sums  $x \oplus x', x \in \mathcal{H}, x' \in \mathcal{H}'$ , and define a hermitian form by  $(x \oplus x' | y \oplus y') = (x | x') + (y | y')$ .  $\mathcal{H}''$  is easily seen to be a Hilbert space, the direct sum of  $\mathcal{H}$  and  $\mathcal{H}'$ .  $\mathcal{H}$  and  $\mathcal{H}'$  can be identified with subspaces of  $\mathcal{H}''$ .*

In general, we shall write  $\mathcal{H}'' = \mathcal{H} \oplus \mathcal{H}'$  to speak of an isomorphic situation :  $\mathcal{H}$  and  $\mathcal{H}'$  are closed subspaces of  $\mathcal{H}$  (hence Hilbert spaces),  $(x | y) = 0$  for  $x \in \mathcal{H}, x' \in \mathcal{H}'$ , and every vector  $x'' \in \mathcal{H}''$  can (uniquely) be written  $x'' = x + x'$  for some  $x \in \mathcal{H}$  and  $x' \in \mathcal{H}'$ . Observe that, given  $\mathcal{H}, \mathcal{H}'$  is uniquely determined :  $\mathcal{H}' = \{x'; \forall x \in \mathcal{H} \quad (x | x') = 0\}$ .

**Definition 22**

If  $\mathcal{H}$  and  $\mathcal{H}'$  are Hilbert spaces, then we can form a new space  $\mathcal{H}'' = \mathcal{H} \otimes \mathcal{H}'$  by considering the vector space of all finite linear combinations (with coefficients in  $\mathcal{C}$ ) of formal expressions  $x \otimes x'$  (with  $x \in \mathcal{H}, x' \in \mathcal{H}'$ ). If one defines  $(\sum_i \alpha_i x_i \otimes x'_i \mid \sum_j \beta_j y_j \otimes y'_j) = \sum_{ij} \alpha_i \beta_j (x_i \mid y_j)(x'_i \mid y'_j)$ , then it is easily shown that this is actually a hermitian form.  $\mathcal{H}''$  is obtained from this prehilbertian space by separation and completion.

Observe that separation amounts to quotient the vector space by the space of vector with a null semi-norm, which is exactly the space generated by the following vectors :

- $x \otimes (x' + y') - x \otimes x' - x \otimes y'$
- $(x + y) \otimes x' - x \otimes x' - y \otimes x'$
- $(\alpha x) \otimes x' - x \otimes (\alpha x')$
- $(\alpha x) \otimes x' - \alpha(x \otimes x')$

**5.2 Bounded operators****Definition 23**

If  $\mathcal{H}$  and  $\mathcal{H}'$  are Hilbert spaces, then a map  $u$  from  $\mathcal{H}$  to  $\mathcal{H}'$  is said to be a bounded operator when the following hold :

- $u$  is linear, i.e.  $u(\alpha x + \beta y) = \alpha u(x) + \beta u(y)$  for  $\alpha, \beta \in \mathcal{C}$  and  $x, y \in \mathcal{H}$ .
- The quantity  $\|u\| = \sup_{\|x\| \leq 1} \|u(x)\|$  is finite ;  $\|u\|$  is the norm of  $u$ .

If  $u$  and  $v$  are bounded operators from  $\mathcal{H}$  to  $\mathcal{H}'$ , if  $\alpha \in \mathcal{C}$ , then one can define bounded operators  $u + v$  and  $\alpha u$  from  $\mathcal{H}$  to  $\mathcal{H}'$ , by means of  $(u + v)(x) = u(x) + v(x)$ ,  $(\alpha u)(x) = \alpha u(x)$ . If  $u$  and  $v$  are bounded operators from  $\mathcal{H}'$  to  $\mathcal{H}''$  and  $\mathcal{H}$  to  $\mathcal{H}'$  respectively, then one can define a bounded operator  $uv$  from  $\mathcal{H}$  to  $\mathcal{H}''$ , by means of  $(uv)(x) = u(v(x))$ . Observe that  $\|u + v\| \leq \|u\| + \|v\|$ ,  $\|\alpha u\| \leq |\alpha| \|u\|$ ,  $\|uv\| \leq \|u\| \|v\|$ . The operators 0 and 1 (the null operator and the identity) have respective norms 0 and 1.

**Proposition 6**

Assume that  $u_i, i = 1, 2$  are bounded operators from  $\mathcal{H}_i$  to  $\mathcal{H}'_i$  ; then

- There is a unique bounded operator  $u_1 \oplus u_2$  from  $\mathcal{H}_1 \oplus \mathcal{H}_2$  to  $\mathcal{H}'_1 \oplus \mathcal{H}'_2$  such that  $(u_1 \oplus u_2)(x_1 \oplus x_2) = u_1(x_1) \oplus u_2(x_2)$
- There is a unique bounded operator  $u_1 \otimes u_2$  from  $\mathcal{H}_1 \otimes \mathcal{H}_2$  to  $\mathcal{H}'_1 \otimes \mathcal{H}'_2$  such that  $(u_1 \otimes u_2)(x_1 \otimes x_2) = u_1(x_1) \otimes u_2(x_2)$

**Definition 24**

A bounded operator  $u$  from  $\mathcal{H}$  to  $\mathcal{H}'$  is said to be an isometry of  $\mathcal{H}$  into  $\mathcal{H}'$  when it preserves the norm, i.e.  $(u(x) | u(x)) = (x | x)$  for all  $x \in \mathcal{H}$ ; this condition is easily seen to imply the more general condition  $(u(x) | u(y)) = (x | y)$ . Among typical isometries, let us mention rotations (in the Hilbert space  $\mathcal{C}^n$ ) and the maps which identify  $\mathcal{H}$  and  $\mathcal{H}'$  with subspaces of  $\mathcal{H} \oplus \mathcal{H}'$ . An isometry has norm 1 (except when the source space is reduced to the null vector). A surjective isometry from  $\mathcal{H}$  onto  $\mathcal{H}'$  turns out to be an isomorphism of structures.

**Example 8**

Assume that  $f$  is a partial injective map from a subset of  $I$  into  $J$ ; then one can define a bounded operator  $u_f$  from  $\ell^2(I)$  into  $\ell^2(J)$  by  $u_f((x_i)) = (y_j)$ , where the sequence  $(y_j)$  is defined by  $y_{f(i)} = x_i$ ,  $y_j = 0$  when  $j \notin \text{rg}(f)$ . (With friendlier notations :

$u_f(\sum \lambda_i \cdot i) = \sum \lambda_i \cdot f(i)$ ). The norm of  $u_f$  is equal to 1 (except when the domain of  $f$  is empty); furthermore in case  $f$  is total, then  $u_f$  is an isometry of  $\ell^2(I)$  into  $\ell^2(J)$ , and in case  $f$  is also surjective, then the isometry is onto. Observe that, if  $f, g$  are partial injections from respectively a subset of  $J$  into  $K$  and a subset of  $I$  into  $J$ , then  $u_{fg} = u_f u_g$ . Similarly if the graph of the partial function  $f$  is the union of the graphs of the partial functions  $g$  and  $h$ , then  $u_f = u_g + u_h$ .

### 5.3 The adjoint of an operator

The main elementary property of the Hilbert space is the following :

**Proposition 7**

The dual  $\tilde{\mathcal{H}}$  of  $\mathcal{H}$  is (semi)-isomorphic to  $\mathcal{H}$ .

PROOF: We first explain the meaning of the result : if  $a \in \mathcal{H}$ , then the map  $\varphi_a$  from  $\mathcal{H}$  to  $\mathcal{C}$  defined by  $\varphi_a(x) = (x | a)$  is linear and continuous (use Cauchy-Schwarz), i.e. is a member of the topological dual of  $\mathcal{H}$ . Conversely, any element of  $\tilde{\mathcal{H}}$  is indeed of the form  $\varphi_a$ . Therefore the map which sends  $a$  to  $\varphi_a$  is a bijection of  $\mathcal{H}$  onto  $\tilde{\mathcal{H}}$ . But this map is not quite linear, since  $\varphi_{\alpha a + \beta b} = \bar{\alpha} \varphi_a + \bar{\beta} \varphi_b$ , i.e. it is linear up to complex conjugation, this is why it is styled "semi-linear". The proposition is established as follows : if  $\varphi$  is a nonzero continuous form on  $\mathcal{H}$ , then one can show (using the median identity and the completeness of the space) that the set  $\{x; \varphi(x) = 1\}$  has exactly one element  $a$  of minimum norm. Then one easily shows that  $\varphi = \varphi_b$ , with  $b = a/(a | a)$ .  $\square$

If  $u$  is a bounded operator from  $\mathcal{H}$  to  $\mathcal{H}'$ , then  $u$  induces a linear map  $\tilde{u}$  from the dual  $\tilde{\mathcal{H}'}$  to the dual  $\tilde{\mathcal{H}}$ , by means of  $\tilde{u}(\varphi) = \varphi \circ u$ . By proposition 7,  $\tilde{u}$  induces in turn a map  $u^*$  from  $\mathcal{H}'$  to  $\mathcal{H}$  defined by :  $\varphi_a \circ u = \varphi_{u^*(a)}$ , in other terms :

**Definition 25**

If  $u$  is a bounded operator from  $\mathcal{H}$  to  $\mathcal{H}'$ , then we define the adjoint  $u^*$  of  $u$ , a map from  $\mathcal{H}'$  to  $\mathcal{H}$  by :  $(u(a) \mid b) = (a \mid u^*(b))$  for all  $a, b \in \mathcal{H}$ .

**Example 9**

If  $f$  is a partial injective function from the subset  $X$  of  $I$  onto the subset  $Y$  of  $J$ , let  $g$  be its inverse ; then with the notations of example 8, we get  $u_f^* = u_g$ .

**Proposition 8**

The adjoint of a bounded operator is still a bounded operator ; furthermore the following hold :

- $(\alpha u + \beta v)^* = \bar{\alpha} u^* + \bar{\beta} v^*$
- $1^* = 1$
- $(uv)^* = v^* u^*$
- $u^{**} = u$
- $\|u^*\| = \|u\|$
- $\|uu^*\| = \|u\|^2$

PROOF: Most of the properties are immediate ; the last one follows from the Cauchy-Schwarz inequality.  $\square$

**Proposition 9**

If  $u$  is an isometry from  $\mathcal{H}$  into  $\mathcal{H}'$ , then  $u^*u = 1$  ; if  $u$  is surjective, then  $uu^* = 1$ .

PROOF:  $(x \mid y) = (u(x) \mid u(y)) = (x \mid u^*u(y))$  ; then  $(x \mid u^*u(y) - y) = 0$  for all  $x \in \mathcal{H}$ , which implies (take  $x = u^*u(y) - y$ )  $u^*u(y) = y$  ; the second half of the proposition is immediate. Observe that we may have  $u^*u = 1$ , but  $uu^* \neq 1$  : consider  $u_f$  when  $f$  is a non-surjective injection of  $I$  into  $J$ .  $\square$

**5.4  $C^*$ -algebras****Definition 26**

A  $C^*$ -algebra  $A$  consists in the following data :

- A complex Banach algebra, with unit ; in particular  $A$  enjoys the properties  $u \neq 0$  implies  $\|u\| \neq 0$ ,  $\|\alpha u\| = |\alpha| \|u\|$ ,  $\|1\| = 1$ ,  $\|u+v\| \leq \|u\| + \|v\|$ ,  $\|uv\| \leq \|u\| \|v\|$ , and  $A$  is complete w.r.t.  $\|\cdot\|$ .

- A unary operation  $(\cdot)^*$ , called the adjunction, or the involution, and which must satisfy the properties of proposition 8.

### Example 10

The algebra  $\mathcal{B}(\mathcal{H})$  of bounded operators from  $\mathcal{H}$  to itself is the most typical  $C^*$ -algebra. But there are other examples, typically the commutative algebra  $\mathcal{C}(X)$  of continuous complex-valued functions on a compact space  $X$  (with pointwise addition and multiplication, the adjunction being complex conjugation etc.). A famous theorem states that any commutative  $C^*$ -algebra is isomorphic to some  $\mathcal{C}(X)$ . The general case is not that simple : however every  $C^*$ -algebra is isomorphic to a subalgebra of some  $\mathcal{B}(\mathcal{H})$ , i.e. it is always possible to represent the elements of a  $C^*$ -algebra as actual bounded operators ; when  $u$  is represented by an element of  $\mathcal{B}(\mathcal{H})$  then we say that  $u$  *operates*, or acts on  $\mathcal{H}$ .

## 5.5 Zoology of operators

$C^*$ -algebras generalize the field of complex numbers ; one can keep this in mind to understand the following zoology of operators :

- An operator  $u$  is *hermitian* when  $u = u^*$  or equivalently  $(u(x) | x)$  is real for all  $x \in \mathcal{H}$ . Hermitian operators clearly generalize real numbers. Among hermitian operators all projections and symmetries (see below), and all operators  $u + u^*$  and  $uu^*$ . By the way,  $v = uu^*$  enjoys a stronger property, namely that  $(v(x) | x) \geq 0$  for  $x \in \mathcal{H}$  : such an operator is said to be a *positive* hermitian operator, and positive hermitians generalize positive reals.
- An operator  $u$  is said to be a *projection* when it is hermitian and  $u^2 = u$  ; in such a case,  $1 - u$  is also a projection. If  $u$  operates on  $\mathcal{H}$ , then the range  $E$  of  $u$  and the range  $F$  of  $1 - u$  are such that  $\mathcal{H} = E \oplus F$ , and  $u$  corresponds to the orthogonal projection of  $\mathcal{H}$  onto the subspace  $E$  : if  $x = e + f$ ,  $e \in E$ ,  $f \in F$ , then  $u(x) = e$ . Projections, which generalize the reals 0, 1 therefore correspond to closed subspaces of the Hilbert space. Non-zero projections have norm 1.
- An operator  $u$  is said to be *unitary* when  $uu^* = u^*u = 1$  ; unitary operators clearly generalize the unit circle. On Hilbert space, being unitary is equivalent to saying that  $(u(x) | u(x)) = (x | x)$ , and that  $u$  is surjective : i.e. unitary operators are represented by isometries of  $\mathcal{H}$ . Their norm is always 1.
- An operator  $u$  is said to be a *symmetry* when it is both hermitian and unitary ; this generalizes the reals 0, 1. If  $u$  is a projection, then  $2u - 1$

is a symmetry, and conversely if  $u$  is a symmetry, then  $(1 + u)/2$  is a projection. Symmetries are indeed yet another way to speak of closed subspaces of  $\mathcal{H}$  : instead of defining a projection from an orthogonal direct sum decomposition, one can introduce the symmetry  $u(e + f) = e - f$ .

- Hermitian and unitary operators share one property, namely that they are normal, i.e. that  $uu^* = u^*u$ . For normal operators, lot of results from finite dimensional algebra can be generalized, typically some forms of diagonalisation. However, one can easily meet non normal operators, especially in geometry of interaction, the typical example being partial isometries.
- $u$  is said to be a *partial isometry* when  $uu^*$  is a projection. This condition indeed implies that  $u^*u$  is a projector (PROOF: consider  $v = uu^*u - u$  ; then  $vv^* = (uu^*)^3 - 2(uu^*)^2 + uu^*$  ; if  $uu^*$  is a projection, then  $vv^* = 0$ , hence  $\|v\|^2 = \|vv^*\| = 0$ , hence  $v = 0$ . From  $uu^*u = u$  one easily gets  $(u^*u)^2 = u^*u$ .  $\square$ ) By symmetry, the conditions  $uu^*$  projection and  $u^*u$  projection are equivalent. On a Hilbert space,  $u$  acts as follows : if  $E$  and  $F$  are the subspaces corresponding to  $u^*u$  and  $uu^*$ , then  $u$  induces an isometry between  $E$  and  $F$ . Nonzero partial isometries have norm 1.
- A partial isometry which is also hermitian is called a *partial symmetry* ; equivalently  $u^* = u$  and  $u^3 = u$ . Symmetries and projections are partial symmetries.

### Example 11

Let us see how the operators  $u_f$  of example 8 react to our zoology : in general, if  $f$  is a partial injection of  $I$  into  $I$ , then  $u_f$  is a partial isometry of  $\ell^2(I)$ .  $u_f$  is normal when the domain and the range of  $f$  coincide.  $u_f$  is unitary when  $f$  is a bijection of  $I$ .  $u_f$  is hermitian (i.e. is a partial symmetry) when  $f$  is an involution of its domain. Finally  $u_f$  is a projection when  $f$  is the identity on its domain.

## 5.6 The algebra $\Lambda^*(L)$

We explain here how  $\lambda^*(L)$  can be completed into a  $C^*$ -algebra. Here we directly refer to the definitions of subsection 2.2. This basically amounts, by topological generalities, to equip  $\lambda^*(L)$  with  $C^*$ -seminorm, and to proceed as usual, separation/completion. A  $C^*$ -seminorm is exactly the same thing as a  $C^*$ -norm, except that it may take the value 0 on nonzero objects, and also that topological completeness is not required. By the way, observe that, if  $R$

is a clause of  $\lambda^*(L)$ , then  $\|R\| = 1$  or  $0$ , (since  $RR^*$  is an idempotent, we get  $\|RR^*\| = \|RR^*\|^2$ , hence  $\|RR^*\| = 1$  or  $0$ , and  $\|R\| = \sqrt{\|RR^*\|} = 1$  or  $0$ ) and therefore  $\|\sum \alpha_i \cdot (P_i \mapsto Q_i)\|$  is bounded by  $\sum |\alpha_i|$  for any  $C^*$ -seminorm. Then the pointwise supremum of any nonempty family of  $C^*$ -seminorms on  $\lambda^*(L)$  is finite, and it is immediate that this supremum is also a  $C^*$ -seminorm ; hence we can define a unique such seminorm on  $\lambda^*(L)$  as soon as there is at least one  $C^*$ -seminorm. For this it obviously suffices to show that  $\lambda^*(L)$  operates on a Hilbert space.

**Proposition 10**

$\lambda^*(L)$  acts on the Hilbert space  $\ell^2(G)$ , where  $G$  is the set of ground propositions of  $L$  in such a way that the sum of  $\lambda^*(L)$  is interpreted by sum of operators, the scalar multiplication of  $\lambda^*(L)$  by scalar multiplication of operators, the product  $\lambda^*(L)$  by composition of operators, and the involution  $*$  of  $\lambda^*(L)$  by adjunction of operators.

PROOF: let  $G$  be the set of ground propositions in  $L$  ; to any clause  $P \mapsto Q$  in  $L$  we can associate an injection  $|P \mapsto Q|$  from a subset of  $G$  into  $G$  :

- If  $g \in G$  unifies with  $Q$ , then the mgu  $\theta$  yields closed values for all the variables in  $Q$ , which are the variables of  $P$ , hence  $P\theta$  is also a ground formula : we set  $|P \mapsto Q|(g) = P\theta$ .
- If  $g \in G$  does not unifies with  $Q$ , then  $|P \mapsto Q|(g)$  is undefined.

It is immediate that  $|P \mapsto Q|$  is a partial injection (this is because the variables of  $Q$  are all present in  $P$ ), and that

$$|P \mapsto Q| \circ |P' \mapsto Q'| = |(P \mapsto Q) \cdot (P' \mapsto Q')|$$

an equation between partial functions that persists when resolution fails, if we interpret  $0$  as the fully undefined function. As in example 8 the partial injections  $|P \mapsto Q|$  induce operators (partial isometries) of the Hilbert space  $\ell^2(G)$  :

$$|P \mapsto Q|(\sum \alpha_i \cdot g_i) = \sum \alpha_i \cdot |P \mapsto Q|(g_i).$$

This immediately extends to all elements of  $\lambda^*(L)$  which are therefore ascribed operators on  $\ell^2(G)$ , and the properties that we have stated immediately follow from definition 8 and example 9.  $\square$

**Definition 27 (The completion)**

$\Lambda^*(L)$  is defined to be the  $C^*$ -algebra obtained by completing  $\lambda^*(L)$  w.r.t. its greatest  $C^*$ -seminorm.

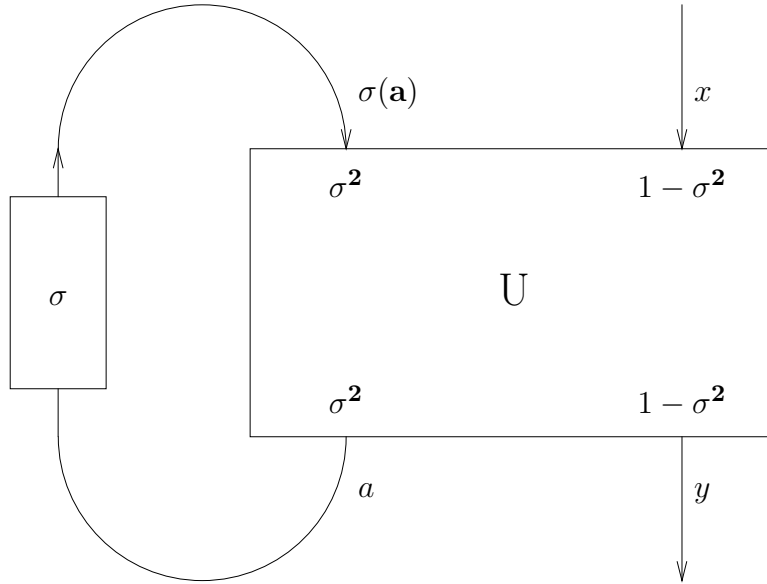
**Remark 17**

It is not difficult to show that a nonzero element of  $\lambda^*(L)$  induces a nonzero operator on  $\ell^2(G)$ , hence the greatest  $C^*$ -seminorm on  $\lambda^*(L)$  is Hausdorff.

## 5.7 The execution formula

The execution formula can also be seen as the solution of a linear equation on a Hilbert space  $\mathcal{H}$  on which  $\Lambda^*(L)$  operates (see subsection 5.4) :

- $U$  produces, given an input  $h \in \mathcal{H}$  an output  $U(h) \in \mathcal{H}$  ;
- $\sigma$  feedbacks certain outputs  $h'$  of  $U$  to inputs  $\sigma(h') \in \mathcal{H}$  ;
- $\sigma^2$  is a projection corresponding to the subspace  $\mathcal{H}'$  on which the feedback is effective ;  $1 - \sigma^2$  is a projection corresponding to the subspace  $\mathcal{H}''$  on which we want to observe the external behavior of the loop. Remark that  $\mathcal{H} = \mathcal{H}' \oplus \mathcal{H}''$ .
- The situation can be summarized by the following figure :



- In other terms, given  $x \in \mathcal{H}''$  we are looking for  $b \in \mathcal{H}$  such that  $(1 - \sigma^2)(b) = x$  (i.e.  $b = c \oplus x$ ) and  $(\sigma U)(b) = b - x$  (i.e.  $c = \sigma U(b)$ ) ; we eventually keep the output value  $y = (1 - \sigma^2)(U(b))$ . In other terms, given  $x \in \mathcal{H}''$  we look for  $a \in \mathcal{H}'$  and  $y \in \mathcal{H}''$  such that

$$U(\sigma(a) \oplus x) = a \oplus y$$

As a linear equation this can be written :

- $(\sigma U)(b) = b - x$ , equivalently  $x = (1 - \sigma U)(b)$  ;



- $b$  is well defined as soon as  $1 - \sigma U$  is invertible, which is the case when  $\sigma U$  is nilpotent ; in that case  $b = (1 - \sigma U)^{-1}(x)$ , which can also be written  $b = (1 - \sigma U)^{-1}(1 - \sigma^2)(x)$  ;
- Therefore  $y = (1 - \sigma^2)U(1 - \sigma U)^{-1}(1 - \sigma^2)(x)$ , i.e.  $y = RES(U, \sigma)(x)$ .

### Remark 18

A naïve way to solve the equation is to write  $U$  and  $\sigma$  as a  $2 \times 2$ -matrices  $(u_{ij})$  and  $(\sigma_{ij})$  (the only nonzero coefficient of  $\sigma_{ij}$  is  $\sigma_{11}$  and  $\sigma_{11}^2 = 1$ ) ; the equation writes as a system :

$$y = u_{21}\sigma_{11}(a) + u_{22}(x) \quad a = u_{11}\sigma_{11}(a) + u_{12}(x)$$

Successive replacements of  $a$  by its value given by the second equation yield :

$$\begin{aligned} y &= u_{21}\sigma_{11}u_{11}\sigma_{11}(a) + u_{21}\sigma_{11}u_{12}(x) + u_{22}(x) = \\ &= u_{21}\sigma_{11}u_{11}\sigma_{11}u_{11}\sigma_{11}(a) + u_{21}\sigma_{11}u_{11}\sigma_{11}u_{12}(x) + u_{21}\sigma_{11}u_{12}(x) + u_{22}(x) \end{aligned}$$

which suggests the infinitary expansion

$$y = u_{11}(x) + \sum_{n \geq 0} u_{21}(\sigma_{11}u_{11})^n \sigma_{11}u_{12}(x)$$

This expansion is by the way legitimate when the sum is finite, which is the case when  $\sigma_{11}u_{11}$  is nilpotent (which is the same as  $\sigma U$  nilpotent). The expression can be rewritten as

$$y = u_{11}(x) + u_{21}(1 - \sigma_{11}u_{11})^{-1}\sigma_{11}u_{12}(x)$$

which is exactly  $RES(U, \sigma)$  (more precisely : its unique non-zero coefficient  $RES(U, \sigma)_{22}$ , which is equal to the coefficient  $EX(U, \sigma)_{22}$ ).

## 5.8 Matrices and direct sums

Among examples of  $cstar$ -algebras, we have the classical example of the algebra  $\mathcal{M}_n(\mathcal{C})$  of  $n \times n$ -matrices with complex coefficients, and more generally  $\mathcal{M}_n(\mathcal{A})$  of  $n \times n$ -matrices with coefficients in a  $C^*$ -algebra  $\mathcal{A}$ . In that case, the adjoint of a matrix  $(a_{ij})$  is the matrix  $(a_{ji}^*)$ . If the algebra  $\mathcal{A}$  acts on a Hilbert space  $\mathcal{H}$ , then  $\mathcal{M}_n(\mathcal{A})$  acts on the  $n$ -ary direct sum  $\mathcal{H} \oplus \dots \oplus \mathcal{H}$  by :

$$(a_{ij})\left(\bigoplus_{i=1}^n x_i\right) = \bigoplus_{i=1}^n \sum_{j=1}^n a_{ij}x_j$$

and this representation is compatible with summation, product, adjunction, etc.

The treatment of geometry of interaction in [G88] involves such matrices, more precisely, the index set consists of the formulas of the concluding sequent (and also of the cut-formulas). This means that in the case of a cut-free proof of  $\vdash A_1, \dots, A_n$ , we see the proof as a linear input/output dependency between  $n$  inputs and  $n$  outputs (one for each  $A_i$ ). Our presentation avoids any mention of matrices ; in fact this is because of the :

**Proposition 11**

$\lambda^*(T^m \cdot n)$  is (isomorphic to) the algebra of  $n \times n$ -matrices with entries in  $\lambda^*(T^m)$ .

PROOF: let us assume for instance  $m = 2$  ; the basic idea is that the matrix  $M_{ij}$  whose entries are all zero but the one of index  $ij$ , which is 1 (i.e. the clause  $p(x, y) \mapsto p(x, y)$ ), is represented by the clause  $p_i(x, y) \mapsto p_j(x, y)$  of  $\lambda^*(T^2 \cdot n)$ .  $\square$

In other terms the predicate letters are used as the indexing of a square matrix. It is funny to observe that 0-ary predicates  $p, q, r$  enjoy

$$(p \mapsto q)(q' \mapsto r) = \delta_{qq'}.(p \mapsto r)$$

with  $\delta_{qq'} = 1$  if  $q = q'$ , 0 otherwise, which makes precise the relation between unification and matrix composition which is behind the previous proposition. When we deal with the rules for the binary connectives, we must in all cases merge two formulas  $A$  and  $B$  into a single formula  $C$  (which is  $A \otimes B$ ,  $A \wp B$ ,  $A \& B$  or  $A \oplus B$ ). This basically amounts to replace a matrix whose indices include  $A$  and  $B$  by another matrix in which these two indices are replaced by a single one,  $C$ . The basic property is that this replacement is a  $*$ -isomorphism, i.e. (we omit the precise definition) it is a map from a  $C^*$ -algebra into another one preserving the structure of  $C^*$ -algebra<sup>9</sup>.

Hence in the case of a binary rule, when we replace the indices  $\Gamma, A, B$  with the indices  $\Gamma, C$ , we are seeking a  $*$ -isomorphism from  $\mathcal{M}_{n+2}(\mathcal{A})$  into  $\mathcal{M}_{n+1}(\mathcal{A})$ . To understand how to construct such a map, imagine that  $u$  operates on a  $n + 2$ -ary direct sum  $\mathcal{H} \oplus \dots \oplus \mathcal{H}$  ; then we would like to see  $u$  as acting on a  $n + 1$ -ary sum. For this it is enough to merge isometrically the last two summands of the  $n + 2$ -ary sum into one summand, by means of an isometry  $\varphi$  of  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$ .

**Proposition 12**

Let  $\varphi(x \oplus y) = P(x) + Q(y)$  be a map from  $\mathcal{H} \oplus \mathcal{H}$  into  $\mathcal{H}$  ; then  $\varphi$  is an isometry iff the following hold :

---

<sup>9</sup>Except perhaps the identity 1 whose image is not requested to be 1 ; the other preservations force it to be a projection.

- $P^*P = Q^*Q = 1$
- $P^*Q = Q^*P = 0$

Furthermore,  $\varphi$  is surjective, i.e. is an isomorphism, when  $PP^* + QQ^* = 1$ .

PROOF: Easy from  $\varphi^*\varphi = 1$  (and  $\varphi\varphi^* = 1$  in case of surjectivity).  $\square$

### Example 12

Consider the  $3 \times 3$ -matrix

$$u = \begin{bmatrix} a_{11} & a_{12} & a_{13} \\ a_{21} & a_{22} & a_{23} \\ a_{31} & a_{32} & a_{33} \end{bmatrix}$$

and let us merge the indices 2, 3 into the index 2 by means of the isometry  $\varphi$  defined from  $P, Q$ ; the result is easily shown to be :

$$v = \begin{bmatrix} a_{11} & a_{12}P^* + a_{13}Q^* \\ Pa_{21} + Qa_{31} & Pa_{22}P^* + Pa_{23}Q^* + Qa_{32}P^* + Qa_{33}Q^* \end{bmatrix}$$

And it is easy to check that this operation on matrices preserves sum, composition, adjunction; the identity matrix is sent to

$$\begin{bmatrix} 1 & 0 \\ 0 & PP^* + QQ^* \end{bmatrix}$$

which is in general only a projection.

Now in order to bridge this with the main text, it is enough to remark that in the algebra  $\lambda^*(T^2)$ , the clauses  $P = p(x, y) \mapsto p(xg, y)$  and  $Q = p(x, y) \mapsto p(xd, y)$  satisfy the conditions of proposition 12. By the way observe that  $PP^* + QQ^* \neq 1$  : this is because there are terms that unify neither with  $xg$  nor with  $xd$ .

## 5.9 Tensorisation and arity

Of course there is another merge that falls into the previous analysis, namely the merge of contexts, in case of a  $\&$ -rule. The only difference is that the pair  $(P', Q')$  chosen works on the second component :  $P' = p(x, y) \mapsto p(x, yg)$  and  $Q' = p(x, y) \mapsto p(x, yd)$ . The fact that  $P', Q'$  coexist can only be explained in terms of tensorization : in general it is possible to define the tensor product  $\mathcal{A} \otimes \mathcal{B}$  of  $C^*$ -algebras  $\mathcal{A}$  and  $\mathcal{B}$ . This new algebra is obtained by completion of the algebra of formal linear combinations of tensors  $a \otimes b$ ,  $a \in \mathcal{A}$ ,  $b \in \mathcal{B}$

with respect to a certain  $C^*$ -norm<sup>10</sup>. Without entering into the technicities of tensor products, the main idea is that, if  $a$  operates on  $\mathcal{H}$  and  $b$  operates on  $\mathcal{H}'$ , then  $a \otimes b$  operates on  $\mathcal{H} \otimes \mathcal{H}'$  by :  $(a \otimes b)(x \otimes y) = a(x) \otimes b(y)$ .

Our way to cope with tensorisation of algebras  $\Lambda^*(L)$  is to play with the arities :

**Proposition 13**

$\lambda^*(T^m)$  contains an isomorphic copy of the  $m$ -ary tensor power of  $\lambda^*(T)$ .

PROOF: typically if  $m = 2$ ,  $p$  is unary and  $q$  is binary, then we can define an isomorphism  $\phi$  from  $\lambda^*(T) \otimes \lambda^*(T)$  into  $\lambda^*(T^2)$ , by

$$\phi((p(t) \mapsto p(u)) \otimes (p(t') \mapsto p(u'))) = q(t, t') \mapsto q(u, u')$$

This isomorphism is of course not surjective.  $\square$

What is behind the proposition is that the term  $tt'$  behaves like the tensor product  $t \otimes t'$ , which can be said pedantically as :

**Proposition 14**

Let  $G$  be the set of ground terms of a term language, including the binary function  $\odot$  ; then the map  $\varphi(g \otimes g') = gg'$  extends to an isometry of  $\ell^2(G) \otimes \ell^2(G)$  into  $\ell^2(G)$ .

Proposition 13 enables us to use  $*$ -isomorphisms to replace the tensorization of algebras  $\Lambda^*(L)$  by other algebras  $\lambda^*(L)$ . In the main text this opportunity is mainly used to induce commutations, since one of the basic facts about the tensor product  $\mathcal{A} \otimes \mathcal{B}$  is that  $u \otimes 1$  commutes with  $1 \otimes v$ . For instance, to come back to our discussion of the pairs  $(P, Q)$  and  $(P', Q')$ , we can introduce  $P'' = p(x) \mapsto p(xg)$  and  $Q'' = p(x) \mapsto p(xd)$  in  $\lambda^*(T)$ , and it is immediate that the operators  $P'' \otimes 1, Q'' \otimes 1, 1 \otimes P'', 1 \otimes Q''$  are sent (by the  $*$ -isomorphism of proposition 13) on  $P, Q, P', Q'$  respectively, and this explains their good relative behavior. In the same way, the constructions of  $\otimes_1(u)$  and  $\otimes_2(u)$  of definition 7 involve tensorizations. And, last but not least, the treatment of exponentials strongly depends on the internalization of tensorization.

**Remark 19**

What is the meaning of lamination (definition 8) ? If we make  $W$  operate on  $(\mathcal{H} \oplus \dots \oplus \mathcal{H}) \otimes \mathcal{H}$  and if  $m$  is a message (definition 5), then  $W(1 \otimes m) = (1 \otimes m)W$ , i.e. that  $W$  belongs to the *commutant* of the set of messages of the form  $1 \otimes m$ .

## References

- [A93] **Asperti, A.** : Linear Logic, Comonads and Optimal reductions, to appear in *Fundamenta Informaticae, Special Issue devoted to categories in computer Science*.

---

<sup>10</sup>In general several  $C^*$ -norms are possible ; the choice includes a greatest and a smallest one

- [B92] **Blass, A.** : A game semantics for linear logic, *Annals Pure Appl. Logic* 56, pp. 183-220, 1992.
- [DR93] **Danos, V. & Regnier, L.** : Proof-nets and the Hilbert space : a summary of Girard's geometry of interaction, *this volume*.
- [G86] **Girard, J.-Y.** : Linear Logic, *Theoretical Computer Science* 50.1, pp. 1-102, 1987.
- [G86A] **Girard, J.-Y.** : Multiplicatives, *Rendiconti del Seminario Matematico dell'Università e Policlinico di Torino, special issue on Logic and Computer Science*, pp. 11-33, 1987.
- [G87] **Girard, J.-Y.** : Towards a Geometry of Interaction, *Categories in Computer Science and Logic, Contemporary Mathematics* 92, AMS 1989, pp. 69-108.
- [G88] **Girard, J.-Y.** : Geometry of Interaction I : interpretation of system F, *Proceedings Logic Colloquium '88*, eds. Ferro & al., pp. 221-260, North Holland 1989.
- [G88A] **Girard, J.-Y.** : Geometry of Interaction II : deadlock-free algorithms, *Proceedings of COLOG-88*, eds Martin-Löf & Mints, pp. 76-93, SLNCS 417.
- [G94] **Girard, J.-Y.** : Proof-nets for additives, *manuscript*, 1994.
- [GAL92] **Gonthier, G. & Abadi, M. & Levy, J.-J.** : The Geometry of Optimal Lambda Reduction, *Proc. 19-th Annual ACM Symposium on Principles of Programming Languages, Albuquerque, New Mexico, ACM Press, New York*, 1992.
- [L93] **Lafont, Y.** : From proof-nets to interaction nets, *this volume*.
- [LMSS90] **Lincoln, P. & Mitchell, J. & Scedrov, A. & Shankar N.** : Decision Problems in Propositional Linear Logic, *Annals of Pure and Applied Logic* 56, 1992, pp. 239-311.
- [LW92] **Lincoln, P. & Winkler, T.** : Constant-only Multiplicative Linear Logic is NP-complete, *Theoretical Computer Science* 135, 1994, pp. 155-159.
- [MR91] **Malacaria, P. & Regnier, L.** : Some results on the interpretation of  $\lambda$ -calculus in operator algebras, *Proc. of LICS '91, IEEE Computer Society Press*, 1991, pp. 63-72.

- [MM93] **Martini, S. & Masini A. :** On the fine structure of the exponential rules, *this volume*.