

LINEAR LOGIC : ITS SYNTAX AND SEMANTICS

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

1 THE SYNTAX OF LINEAR LOGIC

1.1 The connectives of linear logic

Linear logic is not an alternative logic ; it should rather be seen as an extension of usual logic. Since there is no hope to modify the extant classical or intuitionistic connectives¹, linear logic introduces new connectives.

1.1.1 Exponentials : actions vs situations

Classical and intuitionistic logics deal with stable truths :

if A and $A \Rightarrow B$, then B , *but A still holds*.

This is perfect in mathematics, but wrong in real life, since real implication is *causal*. A causal implication cannot be iterated since the conditions are modified after its use ; this process of modification of the premises (conditions) is known in physics as *reaction*. For instance, if A is to spend \$1 on a pack of cigarettes and B is to get them, you lose \$1 in this process, and you cannot do it a second time. The reaction here was that \$1 went out of your pocket. The first objection to that view is that there are in mathematics, in real life, cases where reaction does not exist or can be neglected : think of a lemma which is forever true, or of a Mr. Soros, who has almost an infinite amount of dollars.

1. Witness the fate of non-monotonic “logics” who tried to tamper with logical rules without changing the basic operations ...

Such cases are *situations* in the sense of stable truths. Our logical refinements should not prevent us to cope with situations, and there will be a specific kind of connectives (*exponentials*, “!” and “?”) which shall express the iterability of an action, i.e. the absence of any reaction ; typically $!A$ means to spend as many dollars as one needs. If we use the symbol $\multimap$ (*linear implication*) for causal implication, a usual intuitionistic implication $A \Rightarrow B$ therefore appears as

$$A \Rightarrow B = (!A) \multimap B$$

i.e. A implies B exactly when B is caused by some iteration of A . This formula is the essential ingredient of a faithful translation of intuitionistic logic into linear logic ; of course classical logic is also faithfully translatable into linear logic², so nothing will be lost ... It remains to see what is gained.

1.1.2 The two conjunctions

In linear logic, two conjunctions $\otimes$ (*times*) and $\&$ (*with*) coexist. They correspond to two radically different uses of the word “and”. Both conjunctions express the availability of two actions ; but in the case of $\otimes$, both will be done, whereas in the case of $\&$, only one of them will be performed (but we shall decide which one). To understand the distinction consider A, B, C :

A : to spend \$1,
 B : to get a pack of Camels,
 C : to get a pack of Marlboro.

An action of type A will be a way of taking \$1 out of one’s pocket (there may be several actions of this type since we own several notes). Similarly, there are several packs of Camels at the dealer’s, hence there are several actions of type B . An action type $A \multimap B$ is a way of replacing any specific dollar by a specific pack of Camels.

Now, given an action of type $A \multimap B$ and an action of type $A \multimap C$, there will be no way of forming an action of type $A \multimap B \otimes C$, since for \$1 you will never get what costs \$2 (there will be an action of type $A \otimes A \multimap B \otimes C$, namely getting two packs for \$2). However, there will be an action of type $A \multimap B \& C$, namely the superimposition of both actions. In order to perform this action, we have first to choose which among the two possible actions we want to perform, and then to do the one selected. This is an exact analogue of the computer instruction **if ... then ... else ...** : in this familiar case, the parts **then ...** and **else ...** are available, but only one of them will be done. Although “ $\&$ ” has obvious disjunctive features, it would be technically wrong to view it as a disjunction : the formulas $A \& B \multimap A$ and $A \& B \multimap B$ are both provable (in the

2. With some problems, see 2.2.7

same way “ $\wp$ ”, to be introduced below, is technically a disjunction, but has prominent conjunctive features). There is a very important property, namely the equivalence³ between $!(A \& B)$ and $!A \otimes !B$.

By the way, there are two disjunctions in linear logic :

- “ $\oplus$ ” (*plus*) which is the dual of “ $\&$ ”, expresses the choice of one action between two possible types ; typically an action of type $A \multimap B \oplus C$ will be to get one pack of Marlboro for the dollar, another one is to get the pack of Camels. In that case, we can no longer decide which brand of cigarettes we shall get. In terms of computer science, the distinction $\&/\oplus$ corresponds to the distinction outer/inner non determinism.
- “ $\wp$ ” (*par*) which is the dual of “ $\otimes$ ”, expresses a dependency between two types of actions ; the meaning of $\wp$ is not that easy, let us just say — anticipating on the introduction of linear negation — that $A \wp B$ can either be read as $A^\perp \multimap B$ or as $B^\perp \multimap A$, i.e. “ $\wp$ ” is a symmetric form of “ $\multimap$ ” ; in some sense, “ $\wp$ ” is the constructive contents of classical disjunction.

1.1.3 Linear negation

The most important linear connective is linear negation $(\cdot)^\perp$ (*nil*). Since linear implication will eventually be rewritten as $A^\perp \wp B$, “nil” is the only negative operation of logic. Linear negation behaves like transposition in linear algebra ($A \multimap B$ will be the same as $B^\perp \multimap A^\perp$), i.e. it expresses a *duality*, that is a change of standpoint :

$$\text{action of type } A = \text{reaction of type } A^\perp$$

(other aspects of this duality are output/input, or answer/question).

The main property of $(\cdot)^\perp$ is that $A^{\perp\perp}$ can, without any problem, be identified with A like in classical logic. But (as we shall see in Section 2) linear logic has a very simple *constructive* meaning, whereas the constructive contents of classical logic (which exists, see 2.2.7) is by no means ... obvious. The involutive character of “nil” ensures De Morgan-like laws for all connectives and quantifiers, e.g.

$$\exists x A = (\forall x A^\perp)^\perp$$

which may look surprising at first sight, especially if we keep in mind that the existential quantifier of linear logic is *effective* : typically, if one proves $\exists x A$, then one proves $A[t/x]$ for a certain term t . This exceptional behaviour of “nil” comes from the fact that $A^\perp$ negates (i.e. reacts to) a single action of type A , whereas usual negation only negates some (unspecified) iteration of A , what

3. This is much more than an equivalence, this is a denotational isomorphism, see 2.2.5

usually leads to a Herbrand disjunction of unspecified length, whereas the idea of linear negation is not connected to anything like a Herbrand disjunction. Linear negation is therefore more primitive, but also stronger (i.e. more difficult to prove) than usual negation.

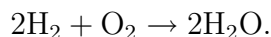
1.1.4 States and transitions

A typical consequence of the excessive focusing of logicians on mathematics is that the notion of *state* of a system has been overlooked.

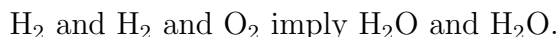
We shall consider below the example of states in (summary !) chemistry, consisting of lists of molecules involved in a reaction (but a similar argumentation can be applied to Petri nets, as first observed by Asperti [4], — a state being a distribution of tokens — or the game of chess — a state being the current position during a play — etc.)

Observe that summary chemistry is modelled according to precise protocols, hence can be formalized : it can eventually be written in mathematics. But in all cases, one will have to introduce an extraneous temporal parameter, and the formalization will explain, in classical logic, how to pass from the state S (modelled as (S, t)) to a new one (modelled as $(S', t + 1)$). This is very awkward, and it would be preferable to ignore this *ad hoc* temporal parameter.

In fact, one would like to represent states by formulas, and transitions by means of implications of states, in such a way that S' is accessible from S exactly when $S \multimap S'$ is provable from the transitions, taken as axioms. But here we meet the problem that, with usual logic, the phenomenon of *updating* cannot be represented. For instance take the chemical equation



A paraphrase of it in current language could be



Common sense knows how to manipulate this as a logical inference ; but this common sense knows that the sense of “and” here is not idempotent (because the proportions are crucial) and that once the starting state has been used to produce the final one, it cannot be reused. The features which are needed here are those of “ $\otimes$ ” to represent “and” and “ $\multimap$ ” to represent “imply” ; a correct representation will therefore be



and it turns out that if we take chemical equations written in this way as axioms, then the notion of linear consequence will correspond to the notion of

accessible state from an initial one. In this example we see that it is crucial that the two following principles of classical logic

$$\begin{aligned} A \wedge B &\Rightarrow A && (\textit{weakening}) \\ A &\Rightarrow A \wedge A && (\textit{contraction}) \end{aligned}$$

become wrong when $\Rightarrow$ is replaced by $\multimap$ and $\wedge$ is replaced by $\otimes$ (contraction would say that the proportions do not matter, whereas weakening would enable us to add an atom of carbon to the left, that would not be present on the right).

To sum up our discussion about states and transitions : the familiar notion of theory — classical logic + axioms — should therefore be replaced by :

$$\textit{theory} = \textit{linear logic} + \textit{axioms} + \textit{current state}.$$

The axioms are there forever ; but the current state is available for a single use : hence once it has been used to prove another state, then the theory is updated, i.e. this other state becomes the next current state. The axioms can of course be replaced by formulas $!A$.

This remark is the basis for potential applications to AI, see [11], this volume : in linear logic certain informations can be logically erased, i.e. the process of revision can be performed by means of logical consequence. What makes it possible is the distinction between formulas $!A$ that speak of stable facts (like the rule of a game) and ordinary ones (that speak about the current state). The impossibility of doing the same thing in classical logic comes from the fact that this distinction makes no sense classically, so any solution to the updating of states would *ipso facto* also be a solution to the updating of the rule of the game⁴.

These dynamical features have been fully exploited in Linear Logic Programming, as first observed in [3]. The basic idea is that the resolution method for linear logic (i.e. proof-search in linear sequent calculus) updates the context, in sharp contrast to intuitionistic proof-search, for which the contexts are monotonic. Updating, inheritance, parallelism are the main features of linear logic programming.

1.1.5 The expressive power of linear logic

Due to the presence of exponentials, linear logic is as expressive as classical or intuitionistic logic. In fact it is more expressive. Here we must be cautious :

4. In particular it would update classical mathematics : can anybody with a mathematical background imagine a minute that commutative algebra can be updated into non-commutative algebra ?

in some sense everything that can be expressed can be expressed in classical logic ... so what ? In fact we have the same problem with intuitionistic logic, which is also “more expressive” than classical logic.

The basic point is that linear logic connectives can express features that classical logic could only handle through complex and *ad hoc* translations. Typically the update of the position m of a pawn inside a chess game with current board M into m' (yielding a new current board M') can be classically handled by means of an implication involving M and M' (and additional features, like temporal markers), whereas the linear implication $m \multimap m'$ will do exactly the same job. The introduction of new connectives is therefore the key to a more manageable way of formalizing ; also the restriction to various fragments opens the area of languages with specific expressive power, e.g. with a given computational complexity.

It is in fact surprising how easily various kinds of abstract machines (besides the pioneering case of Petri nets) can be faithfully translated in linear logic. This is perhaps the most remarkable feature in the study of the complexity of various fragments of linear logic initiated in [25]. See [24], this volume. It is to be remarked that these theorems strongly rely on cut-elimination.

1.1.6 A Far West : non-commutative linear logic

In summary chemistry, all the molecules which contribute to a given state are simultaneously available ; however one finds other kinds of problems in which this is not the case. Typically think of a *stack* $a_0 \dots a_n$ in which a_{n-1} is “hidden” by a_n : if we represent such a state by a conjunction then another classical principle, namely

$$A \wedge B \Rightarrow B \wedge A \quad (\text{exchange})$$

fails, which suggests yet a more drastic modification, i.e. *non-commutative linear logic*. By the way there is an interesting prefiguration of linear logic in the literature, namely Lambek’s *syntactic calculus*, introduced in 1958 to cope with certain questions of linguistic, see [23], this volume. This system is based on a non-commutative $\otimes$, which in turn induces two linear implications. There would be no problems to enrich the system with additives $\&$ and $\oplus$, but the expressive power remains extremely limited. The missing items are exponentials and negation :

- Exponentials stumble on the question of the equivalence between $!(A \& B)$ and $!A \otimes !B$, which is one of the main highway of linear logic : since $\&$ is commutative, the “Times” should be commutative in this case ... or should

one have simultaneously a commutative “Times”, in which case the relation between both types of conjunctions should be understood.

- Linear negation is delicate, since there are several possibilities, e.g. a single negation, like in *cyclic linear logic* as expounded in [27] or two negations, like the two linear implications, in which case the situation may become extremely intricate. Abrusci, see [2], this volume, proposed an interesting solution with two negations.

The problem of finding “the” non-commutative system is delicate, since although many people will agree that non-commutativity makes sense, non-trivial semantics of non-commutativity are not manyfold. In particular a convincing denotational semantics should be set up. By the way, it has been observed from the beginning that non-commutative proof-nets should be planar, which suggests either a planarity restriction or the introduction of braids. Besides the introduction of a natural semantics, the methodology for acknowledging a non-commutative system would also include the gain of expressive power w.r.t. the commutative case.

1.2 Linear sequent calculus

1.2.1 Structural rules

In 1934 Gentzen introduced *sequent calculus*, which is a basic synthetic tool for studying the laws of logic. This calculus is not always convenient to build proofs, but it is essential to study their properties. (In the same way, Hamilton’s equations in mechanics are not very useful to solve practical problems of motion, but they play an essential role when we want to discuss the very principles of mechanics.) Technically speaking, Gentzen introduced *sequents*, i.e. expressions $\Gamma \vdash \Delta$ where $\Gamma (= A_1, \dots, A_n)$ and $\Delta (= B_1, \dots, B_m)$ are finite sequences of formulas. The intended meaning of $\Gamma \vdash \Delta$ is that

$$A_1 \text{ and } \dots \text{ and } A_n \text{ imply } B_1 \text{ or } \dots \text{ or } B_m$$

but the sense of “and”, “imply”, “or” has to be clarified. The calculus is divided into three groups of rules (identity, structural, logical), among which the structural block has been systematically overlooked. In fact, a close inspection shows that the actual meaning of the words “and”, “imply”, “or”, is wholly in the structural group and it is not too excessive to say that a logic is essentially a set of structural rules ! The structural rules considered by Gentzen (respectively *weakening*, *contraction*, *exchange*)

$$\begin{array}{c}
\frac{\Gamma \vdash \Delta}{\Gamma \vdash A, \Delta} \qquad \frac{\Gamma \vdash \Delta}{\Gamma, A \vdash \Delta} \\
\\
\frac{\Gamma \vdash A, A, \Delta}{\Gamma \vdash A, \Delta} \qquad \frac{\Gamma, A, A \vdash \Delta}{\Gamma, A \vdash \Delta} \\
\\
\frac{\Gamma \vdash \Delta}{\sigma(\Gamma) \vdash \tau(\Delta)}
\end{array}$$

are the sequent calculus formulation of the three classical principles already met and criticized. Let us detail them.

Weakening. — Weakening opens the door for fake dependencies : from a sequent $\Gamma \vdash \Delta$ we can get another one $\Gamma' \vdash \Delta'$ by extending the sequences Γ , Δ . Typically, it speaks of causes without effect, e.g. spending \$1 to get nothing — not even smoke —; but it is an essential tool in mathematics (from B deduce $A \Rightarrow B$) since it allows us not to use all the hypotheses in a deduction. It will rightly be rejected from linear logic.

Anticipating on linear sequent calculus, we see that the rule says that $\otimes$ is stronger than $\&$:

$$\frac{\frac{A \vdash A}{A, B \vdash A} \quad \frac{B \vdash B}{A, B \vdash B}}{A, B \vdash A \& B} \quad \frac{}{A \otimes B \vdash A \& B}$$

Affine linear logic is the system of linear logic enriched (?) with weakening. There is no much use for this system since the affine implication between A and B can be faithfully mimicked by $1 \& A \multimap B$. Although the system enjoys cut-elimination, it has no obvious denotational semantics, like classical logic.

Contraction. — Contraction is the fingernail of infinity in propositional calculus : it says that what you have, will always keep, no matter how you use it. The rule corresponds to the replacement of $\Gamma \vdash \Delta$ by $\Gamma' \vdash \Delta'$ where Γ' and Δ' come from Γ and Δ by identifying several occurrences of the same formula (on the same side of “ $\vdash$ ”). To convince oneself that the rule is about infinity (and in fact that without it there is no infinite at all in logic), take the formula $I : \forall x \exists y \, x < y$ (together with others saying that $<$ is a strict order). This axiom has only infinite models, and we show this by exhibiting 1, 2, 3, 4, ... distinct elements ; but, if we want to exhibit 27 distinct elements, we

are actually using I 26 times, and without a principle saying that 26 I can be contracted into one, we would never make it ! In other terms infinity does not mean *many*, but *always*. Another infinitary feature of the rule is that it is the only responsible for undecidability⁵ : Gentzen's subformula property yields a decision method for predicate calculus, provided we can bound the length of the sequents involved in a cut-free proof, and this is obviously the case in the absence of contraction.

In linear logic, both contraction and weakening will be forbidden *as structural rules*. But linear logic is not logic without weakening and contraction, since it would be nonsense not to recover them in some way : we have introduced a new interpretation for the basic notions of logic (actions), but we do not want to abolish the old one (situations), and this is why special connectives (exponentials “!” and “?”) will be introduced, with the two missing structurals as their main rules. The main difference is that we now *control* in many cases the use of contraction, which — one should not forget it — means controlling the length of Herbrand disjunctions, of proof-search, normalization procedures, etc.

Whereas the meaning of weakening is the fact that “ $\otimes$ ” is stronger than “ $\&$ ”, contraction means the reverse implication : using contraction we get :

$$\frac{\frac{A \vdash A \quad B \vdash B}{A \& B \vdash A \quad A \& B \vdash B}}{A \& B, A \& B \vdash A \otimes B} \quad \frac{}{A \& B \vdash A \otimes B}$$

It is difficult to find any evidence of such an implication outside classical logic. The problem is that if we accept contraction without accepting weakening too, we arrive at a very confusing system, which would correspond to an imperfect analysis of causality : consider a petrol engine, in which petrol causes the motion ($P \vdash M$) ; weakening would enable us to call any engine a petrol engine (from $\vdash M$ deduce $P \vdash M$), which is only dishonest, but contraction would be miraculous : from $P \vdash M$ we could deduce $P \vdash P \otimes M$, i.e. that the petrol is not consumed in the causality. This is why the attempts of philosophers to build various *relevance logics* out of the only rejection of weakening were never very convincing⁶

5. If we stay first-order : second-order linear logic is undecidable in the absence of exponentials, as recently shown by Lafont (unpublished), see also [24].

6. These systems are now called substructural logics, which is an abuse, since most of the calculi associated have no cut-elimination

Intuitionistic logic accepts contraction (and weakening as well), but only on the left of sequents : this is done in (what can now be seen as) a very hypocritical way, by restricting the sequents to the case where Δ consists of one formula, so that we are never actually in position to apply a right structural rule. So, when we have a cut-free proof of $\vdash A$, the last rule must be logical, and this has immediate consequences, e.g. if A is $\exists y B$ then $B[t]$ has been proved for some t , etc. These features, that just come from the absence of right contraction, will therefore be present in linear logic, in spite of the presence of an involutive negation.

Exchange. — Exchange expresses the commutativity of multiplicatives : we can replace $\Gamma \vdash \Delta$ with $\Gamma' \vdash \Delta'$ where Γ' and Δ' are obtained from Γ and Δ by permutations of their formulas.

1.2.2 Linear sequent calculus

In order to present the calculus, we shall adopt the following notational simplification : formulas are written from literals $p, q, r, p^\perp, q^\perp, r^\perp$, etc., and constants $\mathbf{1}, \perp, \top, \mathbf{0}$ by means of the connectives $\otimes, \wp, \&, \oplus$ (binary), $!$, $?$ (unary), and the quantifiers $\forall x, \exists x$. Negation is *defined* by De Morgan equations, and linear implication is also a defined connective :

$$\begin{array}{ll}
 \mathbf{1}^\perp := \perp & \perp^\perp := \mathbf{1} \\
 \top^\perp := \mathbf{0} & \mathbf{0}^\perp := \top \\
 (p)^\perp := p^\perp & (p^\perp)^\perp := p \\
 (A \otimes B)^\perp := A^\perp \wp B^\perp & (A \wp B)^\perp := A^\perp \otimes B^\perp \\
 (A \& B)^\perp := A^\perp \oplus B^\perp & (A \oplus B)^\perp := A^\perp \& B^\perp \\
 (!A)^\perp := ?A^\perp & (?A)^\perp := !A^\perp \\
 (\forall x A)^\perp := \exists x A^\perp & (\exists x A)^\perp := \forall x A^\perp
 \end{array}$$

$$A \multimap B := A^\perp \wp B$$

The connectives $\otimes, \wp, \multimap$, together with the neutral elements $\mathbf{1}$ (w.r.t. $\otimes$) and $\perp$ (w.r.t. $\wp$) are called *multiplicatives* ; the connectives $\&$ and $\oplus$, together with the neutral elements $\top$ (w.r.t. $\&$) and $\mathbf{0}$ (w.r.t. $\oplus$) are called *additives* ; the connectives $!$ and $?$ are called *exponentials*. The notation has been chosen for its mnemonic virtues : we can remember from the notation that $\otimes$ is multiplicative and conjunctive, with neutral $\mathbf{1}$, $\oplus$ is additive and disjunctive, with neutral $\mathbf{0}$, that $\wp$ is disjunctive with neutral $\perp$, and that $\&$ is conjunctive with neutral $\top$; the distributivity of $\otimes$ over $\oplus$ is also suggested by our notation.

Sequents are right-sided, i.e. of the form $\vdash \Delta$; general sequents $\Gamma \vdash \Delta$ can be mimicked as $\vdash \Gamma^\perp, \Delta$.

Identity / Negation

$$\frac{}{\vdash A, A^\perp} \quad (\text{identity}) \qquad \frac{\vdash \Gamma, A \quad \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta} \quad (\text{cut})$$

Structure

$$\frac{\vdash \Gamma}{\vdash \Gamma'} \quad (\text{exchange : } \Gamma' \text{ is a permutation of } \Gamma)$$

Logic

$$\begin{array}{ll} \frac{}{\vdash 1} \quad (\text{one}) & \frac{\vdash \Gamma}{\vdash \Gamma, \perp} \quad (\text{false}) \\[10pt] \frac{\vdash \Gamma, A \quad \vdash B, \Delta}{\vdash \Gamma, A \otimes B, \Delta} \quad (\text{times}) & \frac{\vdash \Gamma, A, B}{\vdash \Gamma, A \wp B} \quad (\text{par}) \\[10pt] \frac{}{\vdash \Gamma, \top} \quad (\text{true}) & (\text{no rule for zero}) \\[10pt] \frac{\vdash \Gamma, A \quad \vdash \Gamma, B}{\vdash \Gamma, A \& B} \quad (\text{with}) & \frac{\vdash \Gamma, A}{\vdash \Gamma, A \oplus B} \quad (\text{left plus}) \\ & \frac{\vdash \Gamma, B}{\vdash \Gamma, A \oplus B} \quad (\text{right plus}) \\[10pt] \frac{\vdash ?\Gamma, A}{\vdash ?\Gamma, !A} \quad (\text{of course}) & \frac{\vdash \Gamma}{\vdash \Gamma, ?A} \quad (\text{weakening}) \\[10pt] \frac{\vdash \Gamma, A}{\vdash \Gamma, ?A} \quad (\text{dereliction}) & \frac{\vdash \Gamma, ?A, ?A}{\vdash \Gamma, ?A} \quad (\text{contraction}) \\[10pt] \frac{\vdash \Gamma, A}{\vdash \Gamma, \forall x A} \quad (\text{for all : } x \text{ is not free in } \Gamma) & \frac{\vdash \Gamma, A[t/x]}{\vdash \Gamma, \exists x A} \quad (\text{there is}) \end{array}$$

1.2.3 Comments

The rule for “ $\wp$ ” shows that the comma behaves like a hypocritical “ $\wp$ ” (on the left it would behave like “ $\otimes$ ”) ; “and”, “or”, “imply” are therefore read as “ $\otimes$ ”, “ $\wp$ ”, “ $\multimap$ ”.

In a two-sided version the identity rules would be

$$\frac{}{A \vdash A} \qquad \frac{\Gamma \vdash \Delta, A \quad A, \Lambda \vdash \Pi}{\Gamma, \Lambda \vdash \Delta, \Pi}$$

and we therefore see that the ultimate meaning of the identity group (and the only principle of logic beyond criticism) is that “ A is A ” ; in fact the two rules say that A on the left (represented by $A^\perp$ in the right-sided formulation) is stronger (resp. weaker) than A on the right. The meaning of the identity group is to some extent blurred by our right-sided formulation, since the group may also be seen as the negation group.

The logical group must be carefully examined :

- *multiplicatives and additives* : notice the difference between the rule for $\otimes$ and the rule for $\&$: $\otimes$ requires disjoint contexts (which will never be identified unless $?$ is heavily used) whereas $\&$ works with twice the same context. If we see the contexts of A as the price to pay to get A , we recover our informal distinction between the two conjunctions. In a similar way, the two disjunctions are very different, since $\oplus$ requires one among the premises, whereas $\wp$ requires both).
- *exponentials* : $!$ and $?$ are modalities : this means that $!A$ is simultaneously defined on all formulas : the *of course* rule mentions a context with $?\Gamma$, which means that $?\Gamma$ (or $!\Gamma^\perp$) is known. $!A$ indicates the possibility of using A *ad libitum* ; it only indicates a *potentiality*, in the same way that a piece of paper on the slot of a copying machine *can* be copied ... but nobody would identify a copying machine with all the copies it produces ! The rules for the dual (*why not*) are precisely the three basic ways of *actualizing* this potentiality : erasing (*weakening*), making a single copy (*dereliction*), duplicate ... the machine (*contraction*). It is no wonder that the first relation of linear logic to computer science was the relation to memory pointed out by Yves Lafont in [21].
- *quantifiers* : they are not very different from what they are in usual logic, if we except the disturbing fact that $\exists x$ is now the exact dual of $\forall x$. It is important to remark that $\forall x$ is very close to $\&$ (and that $\exists x$ is very close to $\oplus$).

1.3 Proof-nets

1.3.1 The determinism of computation

For classical and intuitionistic logics, we have an essential property, which dates back to Gentzen (1934), known as the *Hauptsatz*, or cut-elimination theorem ; the *Hauptsatz* presumably traces the borderline between *logic* and

the wider notion of *formal system*. It goes without saying that linear logic enjoys cut-elimination ⁷.

Theorem 1

There is an algorithm transforming any proof of a sequent $\vdash \Gamma$ in linear logic into a cut-free proof of the same sequent.

PROOF. — The proof basically follows the usual argument of Gentzen ; but due to our very cautious treatment of structural rules, the proof is in fact much simpler. There is no wonder, since linear logic comes from a proof-theoretical analysis of usual logic ! $\square$

We have now to keep in mind that the *Hauptsatz* — under various disguises, e.g. normalization in λ -calculus — is used as possible theoretical foundation for computation. For instance consider a text editor : it can be seen as a set of general lemmas (the various subroutines about bracketing, the size of pages etc.), that we can apply to a concrete input, let us say a given page that I write from the keyboard ; observe that the number of such inputs is practically infinite and that therefore our lemmas are about the infinite. Now when I feed the program with a concrete input, there is no longer any reference to infinity ... In mathematics, we could content ourselves with something *implicit* like “your input is correct”, whereas we would be mad at a machine which answers “I can do it” to a request. Therefore, the machine does not only check the correctness of the input, it also demonstrates it by exhibiting the final result, which no longer mentions abstractions about the quasi-infinite *potentiality* of all possible pages. Concretely this elimination of infinity is done by systematically making all concrete replacements — in other terms by running the program. But this is exactly what the algorithm of cut-elimination does.

This is why the *structure* of the cut-elimination procedure is essential. And this structure is quite problematic, since we get problems of *permutation of rules*.

Let us give an example : when we meet a configuration

$$\frac{\frac{\vdash \Gamma, A}{\vdash \Gamma', A} \quad (r) \quad \frac{\vdash A^\perp, \Delta}{\vdash A^\perp, \Delta'} \quad (s)}{\vdash \Gamma', \Delta'} \quad (cut)$$

there is no natural way to eliminate this cut, since the unspecified rules (r) and (s) do not act on A or $A^\perp$; then the idea is to forward the cut upwards :

7. A sequent calculus without cut-elimination is like a car without engine

$$\begin{array}{c}
\frac{\vdash \Gamma, A \quad \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta} \text{ (cut)} \\
\frac{\vdash \Gamma, \Delta}{\vdash \Gamma', \Delta} \text{ (r)} \\
\frac{\vdash \Gamma', \Delta}{\vdash \Gamma', \Delta'} \text{ (s)}
\end{array}$$

But, in doing so, we have decided that rule (r) should now be rewritten *before* rule (s), whereas the other choice

$$\begin{array}{c}
\frac{\vdash \Gamma, A \quad \vdash A^\perp, \Delta}{\vdash \Gamma, \Delta} \text{ (cut)} \\
\frac{\vdash \Gamma, \Delta}{\vdash \Gamma, \Delta'} \text{ (s)} \\
\frac{\vdash \Gamma, \Delta'}{\vdash \Gamma', \Delta'} \text{ (r)}
\end{array}$$

would have been legitimate too. The bifurcation starting at this point is usually irreversible : unless (r) or (s) is later erased, there is no way to interchange them. Moreover the problem stated was completely symmetrical w.r.t. left and right, and we can of course arbitrate between the two possibilities by many bureaucratic tricks ; we can decide that *left* is more important than right, but this choice will at some moment conflict with negation (or implication) whose behaviour is precisely to mimic left by right ... Let's be clear : the *taxonomical* devices that force us to write (r) before (s) or (s) before (r) are not more respectable than the alphabetical order in a dictionary. One should try to get rid of them, or at least, ensure that their effect is limited. In fact *denotational semantics*, see chapter 2 is very important in this respect, since the two solutions proposed have the same denotation. In some sense the two answers — although irreversibly different — are *consistent*. This means that if we eliminate cuts in a proof of an intuitionistic disjunction $\vdash A \vee B$ (or a linear disjunction $\vdash A \oplus B$) and eventually get “a proof of A or a proof of B ”, the side (A or B) is not affected by this bifurcation. However, we would like to get better, namely to have a syntax in which such bifurcations do not occur. In intuitionistic logic (at least for the fragment $\Rightarrow, \wedge, \forall$) this can be obtained by replacing sequent calculus by natural deduction. Typically the two proofs just written will get the same associated deduction ... In other terms natural deduction enjoys a *confluence* (or *Church-Rosser*) property : if $\pi \mapsto \pi', \pi''$ then there is π''' such that $\pi', \pi'' \mapsto \pi'''$, i.e. bifurcations are not irreversible.

1.3.2 Limitations of natural deduction

Let us assume that we want to use natural deduction to deal with proofs in linear logic ; then we run into problems.

(1) Natural deduction is not equipped to deal with classical symmetry : several hypotheses and one (distinguished) conclusion. To cope with symmetrical systems one should be able to accept several conclusions at once ... But then one immediately loses the tree-like structure of natural deductions, with its obvious advantage : a well-determined last rule. Hence natural deduction cannot answer the question. However it is still a serious candidate for an intuitionistic version of linear logic ; we shall below only discuss the fragment $(\otimes, \multimap)$, for which there is an obvious natural deduction system :

$$\begin{array}{c}
 \frac{[A] \vdots B}{A \multimap B} \quad (\multimap\text{-intro}) \qquad \frac{A \quad A \multimap B}{B} \quad (\multimap\text{-elim}) \\
 \\
 \frac{A \quad B}{A \otimes B} \quad (\otimes\text{-intro}) \qquad \frac{A \otimes B \quad \begin{array}{c} [A][B] \vdots C \end{array}}{C} \quad (\otimes\text{-elim})
 \end{array}$$

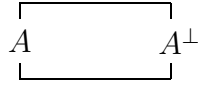
As usual a formula between brackets indicates a *discharge* of hypothesis ; but here the discharge should be *linear*, i.e. exactly one occurrence is discharged (discharging zero occurrence is weakening, discharging two occurrences is contraction). Although this system succeeds in identifying a terrific number of interversion-related proofs, it is not free from serious defects, more precisely :

(2) In the elimination rules the formula which bears the symbol $(\otimes$ or $\multimap)$ is written as a hypothesis ; this is user-friendly, but goes against the actual mathematical structure. Technically this “premise” is in fact the actual conclusion of the rule (think of *main hypotheses* or *headvariables*), which is therefore written upside down. However this criticism is very inessential.

(3) Due to discharge, the introduction rule for $\multimap$ (and the elimination rule for $\otimes$) does not apply to a formula, but to a whole proof. This *global* character of the rule is quite a serious defect.

(4) Last but not least, the elimination rule for $\otimes$ mentions an extraneous formula C which has nothing to do with $A \otimes B$. In intuitionistic natural deduction, we have the same problem with the rules for disjunction and existence which mention an extraneous formula C ; the theory of normalization

which would correspond, in terms of proof-nets to the incestuous configuration :



which is not acknowledged as a proof-net ; in fact in some sense the ultimate meaning of the *correctness* criterion that will be stated below is to forbid such a configuration (and also disconnected ones).

1.3.4 Proof-structures

If we accept the additional links :

$$\frac{A \quad B}{A \otimes B} \quad (\textit{times link}) \qquad \frac{A \quad B}{A \wp B} \quad (\textit{par link})$$

then we can associate to any proof of $\vdash \Gamma$ in linear sequent calculus a graph-like *proof-structure* with as conclusions the formulas of Γ . More precisely :

1. To the identity axiom associate an axiom link.
2. Do not interpret the exchange rule (this rule does not affect conclusions ; however, if we insist on writing a proof-structure on a plane, the effect of the rule can be seen as introducing *crossings* between axiom links ; planar proof-structures will therefore correspond to proofs in some non-commutative variants of linear logic).
3. If a proof-structure β ending with Γ, A and B has been associated to a proof π of $\vdash \Gamma, A, B$ and if one now applies a “par” rule to this proof to get a proof π' of $\vdash \Gamma, A \wp B$, then the structure β' associated to π' will be obtained from β by linking A and B via a *par* link : therefore A and B are no longer conclusions, and a new conclusion $A \wp B$ is created.
4. If π_1 is a proof of $\vdash \Gamma, A$ and π_2 is a proof of $\vdash B, \Delta$ to which proof-structures β_1 and β_2 have been associated, then the proof π' obtained from π_1 and π_2 by means of a times rule is interpreted by means of the proof structure β obtained from β_1 and β_2 by linking A and B together via a *times* link. Therefore A and B are no longer conclusions and a new conclusion $A \otimes B$ is created.
5. If π_1 is a proof of $\vdash \Gamma, A$ and π_2 is a proof of $\vdash A^\perp, \Delta$ to which proof-structures β_1 and β_2 have been associated, then the proof π' obtained from π_1 and π_2 by means of a cut rule is interpreted by means of the proof structure β obtained from β_1 and β_2 by linking A and $A^\perp$ together via a *cut* link. Therefore A and $A^\perp$ are no longer conclusions.

An interesting exercise is to look back at the natural deduction of linear logic and to see how the four rules can be mimicked by proof-structures :

$$\begin{array}{ccc}
 \begin{array}{c} \text{---} A \\ | \\ A^\perp \quad B \\ \hline A^\perp \wp B \end{array} & \begin{array}{c} \text{---} A \quad B^\perp \\ | \quad | \\ A^\perp \quad B^\perp \\ \hline A^\perp \wp B \end{array} & \begin{array}{c} \text{---} A \quad B \\ | \quad | \\ A^\perp \quad B^\perp \\ \hline A^\perp \wp B \end{array} \\
 & \begin{array}{c} A^\perp \wp B \quad A \otimes B^\perp \\ \hline A^\perp \wp B \end{array} & \begin{array}{c} A^\perp \wp B \quad A \otimes B^\perp \\ \hline A^\perp \wp B \end{array} \\
 & & \begin{array}{c} A^\perp \wp B \quad A \otimes B^\perp \\ \hline A^\perp \wp B \end{array}
 \end{array}$$

This shows that — once everything has been put in conclusion —

$$\begin{aligned}
 \neg\circ\text{-intro} &= \otimes\text{-elim} = \text{par link} ; \\
 \neg\circ\text{-elim} &= \otimes\text{-intro} = \text{times link}.
 \end{aligned}$$

1.3.5 Proof-nets

A proof-structure is nothing but a graph whose vertices are (occurrences of) formulas and whose edges are links ; moreover each formula is the conclusion of exactly one link and the premise of at most one link. The formulas which are not premises are the *conclusions* of the structure. Inside proof-structures, let us call *proof-nets* those which can be obtained as the interpretation of sequent calculus proofs. Of course most structures are not nets : typically the definition of a proof-structure does not distinguish between $\otimes$ -links and $\wp$ -links whereas conjunction is surely different from disjunction.

The question which now arises is to find an independent characterization of proof-nets. Let us explain why this is essential :

1. If we define proof-nets from sequent calculus, this means that we work with a proof-structure together with a *sequentialization*, in other terms a step by step construction of this net. But this sequentialization is far from being unique, typically there might be several candidates for the “last rule” of a given proof-net. In practice, we may have a proof-net with a given sequentialization but we may need to use another one : this means that we will spend all of our energy on problems of commutation of rules, as with old

sequent calculus, and we will not benefit too much from the new approach. Typically, if a proof-net ends with a *splitting* $\otimes$ -link, (i.e. a link whose removal induces two disconnected structures), we would like to conclude that the last rule can be chosen as $\otimes$ -rule ; working with a sequentialization this can be proved, but the proof is long and boring, whereas, with a criterion, the result is immediate, since the two components inherit the criterion.

2. The distinction between “and” and “or” has always been explained in semantical terms which ultimately use “and” and “or” ; a purely geometrical characterization would therefore establish the distinction on more intrinsic grounds.

The survey of Yves Lafont [22] (this volume) contains the correctness criterion (first proved in [12] and simplified by Danos and Regnier in [9]) and the sequentialization theorem. From the proof of the theorem, one can extract a quadratic algorithm checking whether or not a given multiplicative proof-structure is a proof-net. Among the uses of multiplicative proof-nets, let us mention the questions of coherence in monoidal closed categories [6].

1.3.6 Cut-elimination for proof-nets

The crucial test for the new syntax is the possibility to handle syntactical manipulations directly at the level of proof-nets (therefore completely ignoring sequent calculus). When we meet a cut link

$$\begin{array}{c} A \quad A^\perp \\ \hline \end{array}$$

we look at links whose conclusions are A and $A^\perp$:

- (1) One of these links is an axiom link, typically :

$$\begin{array}{ccc} & & \vdots \\ & & A^\perp \\ \hline A^\perp & A & A^\perp \\ & & \vdots \\ & & \vdots \end{array}$$

such a configuration can be replaced by

$$\begin{array}{c} \vdots \\ A^\perp \\ \vdots \end{array}$$

however the graphism is misleading, since it cannot be excluded that the two occurrences of $A^\perp$ in the original net are the same ! But this would correspond to a configuration

$$\begin{array}{c} \hline A \quad A^\perp \\ \hline \end{array}$$

in β , and such configurations are excluded by the correctness criterion.

(2) If both formulas are conclusions of logical links for $\otimes$ and $\wp$, typically

$$\frac{\frac{\vdots}{B} \quad \frac{\vdots}{C}}{B \otimes C} \quad \frac{\frac{\vdots}{B^\perp} \quad \frac{\vdots}{C^\perp}}{B^\perp \wp C^\perp}$$

└──────────────────┘

then we can replace it by

$$\frac{\frac{\vdots}{B} \quad \frac{\vdots}{C} \quad \frac{\vdots}{B^\perp} \quad \frac{\vdots}{C^\perp}}{\quad}$$

└──────────────────┘

and it is easily checked that the new structure still enjoys the correctness criterion. This cut-elimination procedure has very nice features :

1. It enjoys a Church-Rosser property (immediate).
2. It is linear in time : simply observe that the proof-net shrinks with any application of steps (1) and (2) ; this linearity is the start of a line of applications to computational complexity.
3. The treatment of the multiplicative fragment is purely local ; in fact all cut-links can be simultaneously eliminated. This must have something to do with parallelism and recently Yves Lafont developed his *interaction nets* as a kind of parallel machine working like proof-nets [22], this volume.

1.3.7 Extension to full linear logic

Proof-nets can be extended to full linear logic. In the case of quantifiers one uses unary links :

$$\frac{A[y/x]}{\forall x A} \quad \frac{A[t/x]}{\exists x A}$$

in the $\forall x$ -link an *eigenvariable* y must be chosen ; each $\forall x$ -link must use a distinct eigenvariable (as the name suggests). The sequentialization theorem can be extended to quantifiers, with an appropriate extension of the correctness criterion.

Additives and neutral elements also get their own notion of proof-nets [17], as well as the part of the exponential rules which does not deal with “!”. Although this extension induces a tremendous simplification of the familiar sequent calculus, it is not as satisfactory as the multiplicative/quantifier case.

Eventually, the only rule which resists to the proof-net technology is the $!$ -rule. For such a rule, one must use a box, see [22]. The box structure has a deep meaning, since the nesting of boxes is ultimately responsible for cut-elimination.

1.4 Is there a unique logic ?

1.4.1 LU

By the turn of the century the situation concerning logic was quite simple : there was basically one logic (classical logic) which could be used (by changing the set of proper axioms) in various situations. Logic was about *pure* reasoning. Brouwer's criticism destroyed this dream of unity : classical logic was not adapted to constructive features and therefore lost its universality. By the end of the century we are now faced with an incredible number of logics — some of them only named “logics” by antiphrasis, some of them introduced on serious grounds—. Is still logic about pure reasoning? In other terms, could there be a way to reunify logical systems —let us say those systems with a good sequent calculus— into a single sequent calculus. Could we handle the (legitimate) distinction classical/intuitionistic not through a change of system, but through a change of formulas? Is it possible to obtain classical effects by restricting one to classical formulas? etc.

Of course there are surely ways to achieve this by cheating, typically by considering a disjoint union of systems . . . All these jokes will be made impossible if we insist on the fact that the various systems represented should freely communicate (and for instance a classical theorem could have an intuitionistic corollary and *vice versa*).

In the unified calculus **LU** see [14], classical, linear, and intuitionistic logics appear as *fragments*. This means that one can define notions of *classical*, *intuitionistic*, or *linear* sequents and prove that a cut-free proof of a sequent in one of these fragments is wholly inside the fragment ; of course a proof with cuts has the right to use arbitrary sequents, i.e. the fragments can freely communicate.

Perhaps the most interesting feature of this new system is that the classical, intuitionistic and linear fragments of **LU** are better behaved than the original sequent calculi. In **LU** the distinction between several styles of maintenance (e.g. “rule of the game” vs “current state”) is particularly satisfactory. But after all, **LU** is little more than a clever reformulation of linear sequent calculus.

1.4.2 LLL and ELL

This dream of unity stumbles on a new fact : the recent discovery [16] of two systems which definitely diverge from classical or intuitionistic logic, **LLL** (*light linear logic*) and **ELL** (elementary linear logic). They come from the basic remark that, in the absence of exponentials, cut-elimination can be performed in linear time. This result (which is conspicuous from a proof-net argument⁸), holds for *lazy* cut-elimination, which does not normalize above $\&$ -rules, and which is enough for algorithmic purposes ; notice that the result holds without regards for the kind of quantifiers —first or second order— used. However the absence of exponentials renders such systems desperately inexpressive. The first attempt to expand this inexpressive system while keeping interesting complexity bounds was not satisfactory : **BLL** (bounded linear logic) [19] had to keep track of polynomial I/O bounds that induced polytime complexity effects, but the price paid was obviously too much.

The problem to solve was therefore to find restriction(s) on the exponentials which ensure :

- cut-elimination, (hence consistency) for naive set-theory, i.e. full comprehension
- the familiar equivalence between $!(A \& B)$ and $!A \otimes !B$

Two systems have been found, both based on the idea that normalization should respect the depth of formulas (with respect to the nesting of $!$ -boxes). Normalization is performed in linear time at depth 0, and induces some duplication of bigger depths, then it is performed at depth 1, etc. and eventually stops, since the total depth does not change. The global complexity therefore depends on the (fixed) global depth and on the number of duplications operated by the “cleansing” of a given level.

- In **LLL** the sizes of inner boxes are multiplied by a factor corresponding to the outer size. The global procedure is therefore done in a time which is a polynomial in the size (with a degree depending of the total depth). Conversely every polytime algorithm can be represented in **LLL**.
- In **ELL** the factor is exponential in the outer size, yielding an elementary complexity for the global procedure, and conversely every elementary algorithm can be represented in **ELL**. **ELL** differs from **LLL** only in the extra principle $!A \otimes !B \multimap !(A \otimes B)$.

LLL may have interesting applications in complexity theory ; **ELL** is expressive enough to accommodate a lot of current mathematics.

8. Remember that the size of a proof-net shrinks during cut-elimination

2 THE SEMANTICS OF LINEAR LOGIC

2.1 The phase semantics of linear logic

The most traditional, and also the less interesting semantics of linear logic associates values to formulas, in the spirit of classical model theory. Therefore it only modelizes provability, and not proofs.

2.1.1 Phase spaces

A *phase space* is a pair $(M, \perp)$, where M is a commutative monoid (usually written multiplicatively) and $\perp$ is a subset of M . Given two subsets X and Y of M , one can define $X \multimap Y := \{m \in M ; \forall n \in X \quad mn \in Y\}$. In particular, we can define for each subset X of M its *orthogonal* $X^\perp := X \multimap \perp$. A *fact* is any subset of M equal to its biorthogonal, or equivalently any subset of the form $Y^\perp$. It is immediate that $X \multimap Y$ is a fact as soon as Y is a fact.

2.1.2 Interpretation of the connectives

The basic idea is to interpret all the operations of linear logic by operations on facts : once this is done the interpretation of the language is more or less immediate. We shall use the same notation for the interpretation, hence for instance $X \otimes Y$ will be the fact interpreting the tensorization of two formulas respectively interpreted by X and Y . This suggests that we already know how to interpret $\perp$, linear implication and linear negation.

1. **times** : $X \otimes Y := \{mn ; m \in X \wedge n \in Y\}^{\perp\perp}$
2. **par** : $X \wp Y := (X^\perp \otimes Y^\perp)^\perp$
3. **1** : $\mathbf{1} := \{1\}^{\perp\perp}$, where 1 is the neutral element of M
4. **plus** : $X \oplus Y := (X \cup Y)^{\perp\perp}$
5. **with** : $X \& Y := X \cap Y$
6. **zero** : $\mathbf{0} := \emptyset^{\perp\perp}$
7. **true** : $\top := M$
8. **of course** : $!X := (X \cap I)^{\perp\perp}$, where I is the set of idempotents of M which belong to $\mathbf{1}$
9. **why not** : $?X := (X^\perp \cap I)^\perp$

(The interpretation of exponentials is an improvement of the original definition of [12] which was awfully *ad hoc*). This is enough to define what is a model of

propositional linear logic. This can easily be extended to yield an interpretation of quantifiers (intersection and biorthogonal of the union). Observe that the definitions satisfy the obvious De Morgan laws relating $\otimes$ and $\wp$ etc. A non-trivial exercise is to prove the associativity of $\otimes$.

2.1.3 Soundness and completeness

It is easily seen that the semantics is sound and complete :

Theorem 2

A formula A of linear logic is provable iff for any interpretation (involving a phase space $(M, \perp)$), the interpretation A^ of A contains the neutral element 1.*

PROOF. — Soundness is proved by a straightforward induction. Completeness involves the building of a specific phase space. In fact we can take as M the monoid of contexts (i.e. multisets of formulas⁹), whose neutral element is the empty context, and we define $\perp := \{\Gamma ; \vdash \Gamma \text{ provable}\}$. If we consider the sets $A^* := \{\Gamma ; \vdash \Gamma, A \text{ provable}\}$, then these sets are easily shown to be facts. More precisely, one can prove (using the identity group) that $A^{\perp*} = A^{*\perp}$. It is then quite easy to prove that in fact A^* is the value of A in a given model : this amounts to prove commutations of the style $(A \otimes B)^* = A^* \otimes B^*$ (these proofs are simplified by the fact that in any De Morgan pair one commutation implies the other, hence we can choose the friendlier commutation). Therefore, if $1 \in A^*$, it follows that $\vdash A$ is provable. $\square$

As far as I know there is no applications for completeness, due to the fact that there is no known concrete phase spaces to which one could restrict and still have completeness. Soundness is slightly more fruitful : for instance Yves Lafont (unpublished, 1994) proved the undecidability of second order propositional linear logic without exponentials by means of a soundness argument. This exploits the fact that a phase semantics is not defined as any algebraic structure enjoying the laws of linear logic, but that it is fully determined from the choice of a commutative monoid and the interpretation $\perp$, as soon as the atoms are interpreted by facts.

2.2 The denotational semantics of linear logic

2.2.1 Implicit versus explicit

First observe that the cut rule is a way to formulate *modus ponens*. It is the essential ingredient of any proof. If I want to prove B , I usually try to prove a useful lemma A and, assuming A , I then prove B . All proofs in nature,

9. We ignore the multiplicity of formulas $?T$, so that I is the set of contexts $?T$

including the most simple ones, are done in this way. Therefore, there is an absolute evidence that the cut rule is the only rule of logic that cannot be removed : without cut it is no longer possible to *reason*.

Now against common sense Gentzen proved his *Hauptsatz* ; for classical and intuitionistic logics (and remember that can be extended to linear logic without problems). This result implies that we can make proofs without cut, i.e. without lemmas (i.e. without modularity, without ideas, etc.). For instance if we take an intuitionistic disjunction $A \vee B$ (or a linear plus $A \oplus B$) then a cut-free proof of it must contain a proof of A or a proof of B . We see at once that this is artificial : who in real life would state $A \vee B$ when he has proved A ? If we want to give a decent status to proof-theory, we have to explain this contradiction.

Formal reasoning (any reasoning) is about *implicit* data. This is because it is more convenient to *forget*. So, when we prove $A \vee B$, we never know which side holds. However, there is —inside the sequent calculus formulation— a completely artificial use of the rules, i.e. to prove without the help of cut ; this artificial subsystem is completely *explicit*. The result of Gentzen is a way to replace a proof by another without cut, which makes explicit the contents of the original proof. Variants of the Gentzen procedure (normalization in natural deduction or in λ -calculus) should also be analysed in that way.

2.2.2 Generalities about denotational semantics

The purpose of *denotational* semantics is precisely to analyse this implicit contents of proofs. The name comes from the old Fregean opposition *sense/denotation* : the denotation is what is implicit in the sense.

The kind of semantics we are interested in is concrete, i.e. to each proof π we associate a set π^* . This map can be seen as a way to define an equivalence $\approx$ between proofs ($\pi \approx \pi'$ iff $\pi^* = \pi'^*$) of the same formulas (or sequents), which should enjoy the following :

1. if π normalizes to π' , then $\pi \approx \pi'$;
2. $\approx$ is non-degenerated, i.e. one can find a formula with at least two non-equivalent proofs ;
3. $\approx$ is a congruence : this means that if π and π' have been obtained from λ and λ' by applying the same logical rule, and if $\lambda \approx \lambda'$, then $\pi \approx \pi'$;
4. certain canonical isomorphisms are satisfied ; among those which are crucial let us mention :

- involutivity of negation (hence De Morgan),
- associativity of “par” (hence of “times”).

Let us comment these points :

- (1) says that $\approx$ is about cut-elimination.
- (2) : of course if all proofs of the same formula are declared to be equivalent, the contents of $\approx$ is empty.
- (3) is the analogue of a Church-Rosser property, and is the key to a modular approach to normalization.
- (4) : another key to modularity is commutation, which means that certain sequences of operations on proofs are equivalent w.r.t. $\approx$. It is clear that the more commutation we get the better, and that we cannot ask too much a priori. However, the two commutations mentioned are a strict minimum without which we would get a mess :
 - involutivity of negation means that we have not to bother about double negations ; in fact this is the semantical justification of our choice of a *defined* negation.
 - associativity of “par” means that the bracketing of a ternary “par” is inessential ; furthermore, associativity renders possible the identification of $A \multimap (B \multimap C)$ with $(A \otimes B) \multimap C$.

The denotational semantics we shall present is a simplification of Scott domains which has been obtained by exploiting the notion of stability due to Berry (see [18] for a discussion). These drastically simplified Scott domains are called *coherent spaces* ; these spaces were first intended as denotational semantics for intuitionistic logic, but it turned out that there were a lot of other operations hanging around. Linear logic first appeared as a kind of linear algebra built on coherent spaces ; then linear sequent calculus was extracted out of the semantics. Recently Ehrhard, see [10], this volume, refined coherent semantics into *hypercoherences*, with applications to the question of sequentiality.

2.2.3 Coherent spaces

Definition 1

A coherent space is a reflexive undirected graph. In other terms it consists of a set $|X|$ of atoms together with a compatibility or coherence relation between atoms, noted $x \circ y$ or $x \circ y \text{ [mod } X]$ if there is any ambiguity as to X .

A clique a in X (notation $a \sqsubset X$) is a subset a of X made of pairwise coherent atoms : $a \sqsubset X$ iff $\forall x \forall y (x \in a \wedge y \in a \Rightarrow x \subset y)$. In fact a coherent space can be also presented as a set of cliques ; when we want to emphasize the underlying graph $(|X|, \subset)$ we call it the web of X .

Besides coherence we can also introduce

- *strict coherence* : $x \frown y$ iff $x \subset y$ and $x \neq y$,
- *incoherence* : $x \asymp y$ iff $^\perp(x \frown y)$,
- *strict incoherence* : $x \smile y$ iff $^\perp(x \subset y)$.

Any of these four relations can serve as a definition of coherent space. Observe fact that $\asymp$ is the negation of $\frown$ and not of $\subset$; this due to the reflexivity of the web.

Definition 2

Given a coherent space X , its linear negation $X^\perp$ is defined by

- $|X^\perp| = |X|$,
- $x \subset y \text{ [mod } X^\perp] \text{ iff } x \asymp y \text{ [mod } X]$.

In other terms, linear negation is nothing but the exchange of coherence and incoherence. It is obvious that linear negation is involutive : $X^{\perp\perp} = X$.

Definition 3

Given two coherent spaces X and Y , the multiplicative connectives $\otimes$, $\wp$, $\multimap$ define a new coherent space Z with $|Z| = |X| \otimes |Y|$; coherence is defined by

- $(x, y) \subset (x', y') \text{ [mod } X \otimes Y] \text{ iff}$
 $x \subset x' \text{ [mod } X] \text{ and } y \subset y' \text{ [mod } Y]$,
- $(x, y) \frown (x', y') \text{ [mod } X \wp Y] \text{ iff}$
 $x \frown x' \text{ [mod } X] \text{ or } y \frown y' \text{ [mod } Y]$,
- $(x, y) \frown (x', y') \text{ [mod } X \multimap Y] \text{ iff}$
 $x \subset x' \text{ [mod } X] \text{ implies } y \frown y' \text{ [mod } Y]$.

Observe that $\otimes$ is defined in terms of $\subset$ but $\wp$ and $\multimap$ in terms of $\frown$. A lot of useful isomorphisms can be obtained

1. De Morgan equalities : $(X \otimes Y)^\perp = X^\perp \wp Y^\perp$; $(X \wp Y)^\perp = X^\perp \otimes Y^\perp$;
 $X \multimap Y = X^\perp \wp Y$;

2. commutativity isomorphisms : $X \otimes Y \simeq Y \otimes X$; $X \wp Y \simeq Y \wp X$;
 $X \multimap Y \simeq Y^\perp \multimap X^\perp$;
3. associativity isomorphisms : $X \otimes (Y \otimes Z) \simeq (X \otimes Y) \otimes Z$; $X \wp (Y \wp Z) \simeq (X \wp Y) \wp Z$;
 $X \multimap (Y \multimap Z) \simeq (X \multimap Y) \multimap Z$;

Definition 4

Up to isomorphism there is a unique coherent space whose web consists of one atom 0, this space is self dual, i.e. equal to its linear negation. However the algebraic isomorphism between this space and its dual is logically meaningless, and we shall, depending on the context, use the notation $\mathbf{1}$ or the notation $\perp$ for this space, with the convention that $\mathbf{1}^\perp = \perp$, $\perp^\perp = \mathbf{1}$.

This space is neutral w.r.t. multiplicatives, namely $X \otimes \mathbf{1} \simeq X$, $X \wp \perp \simeq X$, $\mathbf{1} \multimap X \simeq X$, $X \multimap \perp \simeq X^\perp$.

This notational distinction is mere preciousity ; one of the main drawbacks of denotational semantics is that it interprets logically irrelevant properties ... but nobody is perfect.

Definition 5

Given two coherent spaces X and Y the additive connectives $\&$ and $\oplus$, define a new coherent space Z with $|Z| = |X| + |Y|$ ($= |X| \otimes \{0\} \cup |Y| \otimes \{1\}$) ; coherence is defined by

- $(x, 0) \subset (x', 0)$ [mod Z] iff $x \subset x'$ [mod X],
- $(y, 1) \subset (y', 1)$ [mod Z] iff $y \subset y'$ [mod Y],
- $(x, 0) \frown (y, 1)$ [mod $X \& Y$],
- $(x, 0) \smile (y, 1)$ [mod $X \oplus Y$].

A lot of useful isomorphisms are immediately obtained :

- De Morgan equalities : $(X \& Y)^\perp = X^\perp \oplus Y^\perp$; $(X \oplus Y)^\perp = X^\perp \& Y^\perp$;
- commutativity isomorphisms : $X \& Y \simeq Y \& X$; $X \oplus Y \simeq Y \oplus X$;
- associativity isomorphisms : $X \& (Y \& Z) \simeq (X \& Y) \& Z$; $X \oplus (Y \oplus Z) \simeq (X \oplus Y) \oplus Z$;
- distributivity isomorphisms : $X \otimes (Y \oplus Z) \simeq (X \otimes Y) \oplus (X \otimes Z)$; $X \wp (Y \& Z) \simeq (X \wp Y) \& (X \wp Z)$; $X \multimap (Y \& Z) \simeq (X \multimap Y) \& (X \multimap Z)$;
 $(X \oplus Y) \multimap Z \simeq (X \multimap Z) \& (Y \multimap Z)$.

The other distributivities fail ; for instance $X \otimes (Y \& Z)$ is not isomorphic to $(X \otimes Y) \& (X \otimes Z)$.

Definition 6

There is a unique coherent space with an empty web. Although this space is also self dual, we shall use distinct notations for it and its negation : $\top$ and $\mathbf{0}$.

These spaces are neutral w.r.t. additives : $X \oplus \mathbf{0} \simeq X$, $X \& \top \simeq X$, and absorbing w.r.t. multiplicatives $X \otimes \mathbf{0} \simeq \mathbf{0}$, $X \wp \top \simeq \top$, $\mathbf{0} \multimap X \simeq \top$, $X \multimap \top \simeq \top$.

2.2.4 Interpretation of MALL

MALL is the fragment of linear logic without the exponentials “!” and “?”. In fact we shall content ourselves with the propositional part and omit the quantifiers. If we wanted to treat the quantifiers, the idea would be to essentially interpret $\forall x$ and $\exists x$ respectively “big” $\&$ and $\oplus$ indexed by the domain of interpretation of variables ; the precise definition involves considerable bureaucracy for something completely straightforward. The treatment of second-order quantifiers is of course much more challenging and will not be explained here. See for instance [12].

Once we decided to ignore exponentials and quantifiers, everything is ready to interpret formulas of **MALL** : more precisely, if we assume that the atomic propositions $p, q, r, \dots$ of the language have been interpreted by coherent spaces $p^*, q^*, r^*, \dots$, then any formula A of the language is interpreted by a well-defined coherent space A^* ; moreover this interpretation is consistent with the definitions of linear negation and implication (i.e. $A^{\perp*} = A^{*\perp}$, $(A \multimap B)^* = A^* \multimap B^*$). It remains to interpret sequents ; the idea is to interpret $\vdash \Gamma$ ($= \vdash A_1, \dots, A_n$) as $A_1^* \wp \dots \wp A_n^*$. More precisely

Definition 7

If $\vdash \Xi (= \vdash X_1, \dots, X_n)$ is a formal sequent made of coherent spaces, then the coherent space $\vdash \Xi$ is defined by

1. $|\vdash \Xi| = |X_1| \otimes \dots \otimes |X_n|$; we use the notation $x_1 \dots x_n$ for the atoms of $\vdash \Xi$.

2. $x_1 \dots x_n \frown y_1 \dots y_n \Leftrightarrow \exists i \ x_i \frown y_i$.

If $\vdash \Gamma (= \vdash A_1, \dots, A_n)$ is a sequent of linear logic, then $\vdash \Gamma^*$ will be the coherent space $\vdash A_1^*, \dots, A_n^*$.

The next step is to interpret proofs ; the idea is that a proof π of $\vdash \Gamma$ will be interpreted by a clique $\pi^* \sqsubset \vdash \Gamma^*$. In particular (since sequent calculus is eventually about proofs of singletons $\vdash A$) a proof π of $\vdash A$ is interpreted by a clique of $\vdash A^*$ i.e. a clique in A^* .

Definition 8

1. The identity axiom $\vdash A, A^\perp$ of linear logic is interpreted by the set $\{xx; x \in |A^*|\}$.
2. Assume that the proofs π of $\vdash \Gamma, A$ and λ of $\vdash A^\perp, \Delta$ have been interpreted by cliques π^* and λ^* in the associated coherent spaces; then the proof ρ of $\vdash \Gamma, \Delta$ obtained by means of a cut rule between π and λ is interpreted by the set $\rho^* = \{\underline{x}\underline{x'}; \exists z (\underline{x}z \in \pi^* \wedge z\underline{x'} \in \lambda^*)\}$.
3. Assume that the proof π of $\vdash \Gamma$ has been interpreted by a clique $\pi^* \sqsubset \Vdash \Gamma^*$, and that ρ is obtained from π by an exchange rule (permutation σ of Γ); then ρ^* is obtained from π^* by applying the same permutation $\rho^* = \{\sigma(\underline{x}); \underline{x} \in \pi^*\}$.

All the sets constructed by our definition are cliques; let us remark that in the case of cut, the atom z of the formula is uniquely determined by $\underline{x}$ and $\underline{x'}$.

Definition 9

1. The axiom $\vdash \mathbf{1}$ of linear logic is interpreted by the clique $\{0\}$ of $\mathbf{1}$ (if we call 0 the only atom of $\mathbf{1}$).
2. The axioms $\vdash \Gamma, \top$ of linear logic are interpreted by void cliques (since $\top$ has an empty web, the spaces $(\vdash \Gamma, \top)^*$ have empty webs as well).
3. If the proof ρ of $\vdash \Gamma, \perp$ comes from a proof π of $\vdash \Gamma$ by a falsum rule, then we define $\rho^* = \{\underline{x}0; \underline{x} \in \pi^*\}$.
4. If the proof ρ of $\vdash \Gamma, A \wp B$ comes from a proof π of $\vdash \Gamma, A, B$ by a par rule, we define $\rho^* = \{\underline{x}(y, z); \underline{x}yz \in \pi^*\}$.
5. If the proof ρ of $\vdash \Gamma, A \otimes B, \Delta$ comes from proofs π of $\vdash \Gamma, A$ and λ of $\vdash B, \Delta$ by a times rule, then we define $\rho^* = \{\underline{x}(y, z)\underline{x'}; \underline{x}y \in \pi^* \wedge z\underline{x'} \in \lambda^*\}$.
6. If the proof ρ of $\vdash \Gamma, A \oplus B$ comes from a proof π of $\vdash \Gamma, A$ by a left plus rule, then we define $\rho^* = \{\underline{x}(y, 0); \underline{x}y \in \pi^*\}$; if the proof ρ of $\vdash \Gamma, A \oplus B$ comes from a proof π of $\vdash \Gamma, B$ by a right plus rule, then we define $\rho^* = \{\underline{x}(y, 1); \underline{x}y \in \pi^*\}$.
7. If the proof ρ of $\vdash \Gamma, A \& B$ comes from proofs π of $\vdash \Gamma, A$ and λ of $\vdash \Gamma, B$ by a with rule, then we define $\rho^* = \{\underline{x}(y, 0); \underline{x}y \in \pi^*\} \cup \{\underline{x}(y, 1); \underline{x}y \in \lambda^*\}$.

Observe that (4) is mainly a change of bracketing, i.e. does strictly nothing; if $|A| \cap |B| = \emptyset$ then one can define $A \& B$, $A \oplus B$ as unions, in which case (6) is read as $\rho^* = \pi^*$ in both cases, and (7) is read $\rho^* = \pi^* \cup \lambda^*$.

It is of interest (since this is deeply hidden in Definition 9) to stress the relation between *linear implication* and *linear maps*:

Definition 10

Let X and Y be coherent spaces ; a linear map from X to Y consists in a function F such that

1. if $a \sqsubset X$ then $F(a) \sqsubset Y$,
2. if $\bigcup b_i = a \sqsubset X$ then $F(a) = \bigcup F(b_i)$,
3. if $a \cup b \sqsubset X$, then $F(a \cap b) = F(a) \cap F(b)$.

The last two conditions can be rephrased as the preservation of disjoint unions.

Proposition 1

There is a 1-1 correspondence between linear maps from X to Y and cliques in $X \multimap Y$; more precisely

- to any linear F from X to Y , associate $\text{Tr}(F) \sqsubset X \multimap Y$ (the trace of F)

$$\text{Tr}(F) = \{(x, y) ; y \in F(\{x\})\},$$

- to any $A \sqsubset X \multimap Y$ associate a linear function $A(\cdot)$ from X to Y

$$\text{if } a \sqsubset X, \text{ then } A(a) = \{y ; \exists x \in a (x, y) \in A\}.$$

PROOF. — The proofs that $\text{Tr}(A(\cdot)) = A$ and $\text{Tr}(F)(\cdot) = F$ are left to the reader. In fact the structure of the space $X \multimap Y$ has been obtained so as to get this property and not the other way around. $\square$

2.2.5 Exponentials

Definition 11

Let X be a coherent space ; we define $\mathcal{M}(X)$ to be the free commutative monoid generated by $|X|$. The elements of $\mathcal{M}(X)$ are all the formal expressions $[x_1, \dots, x_n]$ which are finite multisets of elements of $|X|$. This means that $[x_1, \dots, x_n]$ is a sequence in $|X|$ defined up to the order. The difference with a subset of $|X|$ is that repetitions of elements matter. One easily defines the sum of two elements of $\mathcal{M}(X)$:

$[x_1, \dots, x_n] + [y_1, \dots, y_n] = [x_1, \dots, x_n, y_1, \dots, y_n]$, and the sum is generalized to any finite set. The neutral element of $\mathcal{M}(X)$ is written $[\]$.

If X is a coherent space, then $!X$ is defined as follows :

- $!|X| = \{[x_1, \dots, x_n] \in \mathcal{M}(X) ; x_i \subset x_j \text{ for all } i \text{ and } j\}$,
- $\sum[x_i] \subset \sum[y_j] \text{ [mod } !X] \text{ iff } x_i \subset y_j \text{ for all indices } i \text{ and } j$.

If X is a coherent space, then $?X$ is defined as follows :

- $|?X| = \{[x_1, \dots, x_n] \in \mathcal{M}(X) ; x_i \asymp x_j \text{ for all } i \text{ and } j\},$
- $\Sigma[x_i] \frown \Sigma[y_j] \text{ [mod } ?X] \text{ iff } x_i \frown y_j \text{ for some pair of indices } i \text{ and } j.$

Among remarkable isomorphisms let us mention

- De Morgan equalities : $(!X)^\perp = ?(X^\perp) ; (?X)^\perp = !(X^\perp) ;$
- the exponentiation isomorphisms : $!(X \& Y) \simeq (!X) \otimes (!Y) ; ?(X \oplus Y) \simeq (?X) \wp (?Y),$ together with the “particular cases” $!\top \simeq \mathbf{1} ; ?\mathbf{0} \simeq \perp.$

Definition 12

1. Assume that the proof π of $\vdash ?\Gamma, A$ has been interpreted by a clique π^* ; then the proof ρ of $\vdash ?\Gamma, !A$ obtained from π by an of course rule is interpreted by the set

$$\rho^* = \{(\underline{x}_1 + \dots + \underline{x}_k)[a_1, \dots, a_k] ; \underline{x}_1 a_1, \dots, \underline{x}_k a_k \in \pi^*\}.$$

About the notation : if $?\Gamma$ is $?B^1, \dots, ?B^n$ then each $\underline{x}_i$ is $x_i^1, \dots, x_i^n$ so $\underline{x}_1 + \dots + \underline{x}_k$ is the sequence $x_1^1 + \dots + x_k^1, \dots, x_1^n + \dots + x_k^n ; [a_1, \dots, a_k]$ refers to a multiset. What is implicit in the definition (but not obvious) is that we take only those expressions $(\underline{x}_1 + \dots + \underline{x}_k)[a_1, \dots, a_k]$ such that $\underline{x}_1 + \dots + \underline{x}_k \in \Vdash ?\Gamma$ (this forces $[a_1, \dots, a_k] \in !|A|$).

2. Assume that the proof π of $\vdash \Gamma$ has been interpreted by a clique π^* ; then the proof ρ of $\vdash \Gamma, ?A$ obtained from π by a weakening rule is interpreted by the set $\rho^* = \{\underline{x}[\] ; \underline{x} \in \pi^*\}.$
3. Assume that the proof π of $\vdash \Gamma, ?A, ?A$ has been interpreted by a clique π^* ; then the proof ρ of $\vdash \Gamma, ?A$ obtained from π by a contraction rule is interpreted by the set $\rho^* = \{\underline{x}(a + b) ; \underline{x}ab \in \pi^* \wedge a \asymp b\}.$
4. Assume that the proof π of $\vdash \Gamma, A$ has been interpreted by a clique π^* ; then the proof ρ of $\vdash \Gamma, ?A$ obtained from π by a dereliction rule is interpreted by the set $\rho^* = \{\underline{x}[a] ; \underline{x}a \in \pi^*\}.$

2.2.6 The bridge with intuitionism

First the version just given for the exponentials is not the original one, which was using sets instead of multisets. The move to multisets is a consequence of recent progress on classical logic [13] for which this replacement has deep consequences. But as far as linear and intuitionistic logic are concerned, we can work with sets, and this is what will be assumed here. In particular $\mathcal{M}(X)$ will be replaced by the monoid of finite subsets of X , and sum will be replaced by union. The web of $!X$ will be the set X_{fin} of all finite cliques of X .

Definition 13

Let X and Y be coherent spaces ; a stable map from X to Y is a function F such that

1. if $a \sqsubset X$ then $F(a) \sqsubset Y$,
2. assume that $a = \bigcup b_i$, where b_i is directed with respect to inclusion, then
 $F(a) = \bigcup F(b_i)$,
3. if $a \cup b \sqsubset X$, then $F(a \cap b) = F(a) \cap F(b)$.

Definition 14

Let X and Y be coherent spaces ; then we define the coherent space $X \Rightarrow Y$ as follows :

- $|X \Rightarrow Y| = X_{fin} \otimes |Y|$,
- $(a, y) \subset (a', y')$ iff (1) and (2) hold :
 1. $a \cup a' \sqsubset X \Rightarrow y \subset y'$,
 2. $a \cup a' \sqsubset X \wedge a \neq a' \Rightarrow y \frown y'$.

Proposition 2

There is a 1-1 correspondence between stable maps from X to Y and cliques in $X \Rightarrow Y$; more precisely

1. to any stable F from X to Y , associate $\text{Tr}(F) \sqsubset X \Rightarrow Y$ (the trace of F)

$$\text{Tr}(F) = \{(a, y) ; a \sqsubset X \wedge y \in F(a) \wedge \forall a' \subset a (y \in F(a') \Rightarrow a' = a)\}$$

2. to any $A \sqsubset X \Rightarrow Y$, associate a stable function $A(\cdot)$ from X to Y

$$\text{if } a \sqsubset X, \text{ then } A(a) = \{y ; \exists b \subset a ((b, y) \in A)\}.$$

PROOF. — The essential ingredient is the normal form theorem below. □

Theorem 3

Let F be a stable function from X to Y , let $a \sqsubset X$, let $y \in F(a)$; then

1. there exists $a_0 \subset a$, a_0 finite such that $y \in F(a_0)$,
2. if a_0 is chosen minimal w.r.t. inclusion, then it is unique.

PROOF. — (1) follows from $a = \bigcup a_i$, the directed union of its finite subsets ; $z \in F(\bigcup a_i) = \bigcup F(a_i)$ hence $z \in F(a_i)$ for some i .

(2) : given two solutions a_0, a_1 included in a , we get $z \in F(a_0) \cap F(a_1) = F(a_0 \cap a_1)$; if a_0 is minimal w.r.t. inclusion, this forces $a_0 \cap a_1 = a_0$, hence $a_0 \subset a_1$. $\square$

This establishes the basic bridge with linear logic, since $X \Rightarrow Y$ is strictly the same thing as $!X \multimap Y$ (if we use sets instead of multisets). In fact one can translate intuitionistic logic into linear logic as follows :

$$\begin{aligned} p^* &:= p \quad (p \text{ atomic}), \\ (A \Rightarrow B)^* &:= !A^* \multimap B^*, \\ (A \wedge B)^* &:= A^* \& B^*, \\ (\forall x A)^* &:= \forall x A^*, \\ (A \vee B)^* &:= !A^* \oplus !B^*, \\ (\exists x A)^* &:= \exists x !A^*, \\ (\perp A)^* &:= !A^* \multimap \mathbf{0}. \end{aligned}$$

and prove the following result: $\Gamma \vdash A$ is intuitionistically provable iff $!\Gamma^* \vdash A^*$ (i.e. $\vdash ?\Gamma^{*\perp}, A^*$) is linearly provable. The possibility of such a faithful translation is of course a major evidence for linear logic, since it links it with intuitionistic logic in a strong sense. In particular linear logic can *at least* be accepted as a way of analysing intuitionistic logic.

2.2.7 The bridge with classical logic

Let us come back to exponentials ; the space $!X$ is equipped with two maps :

$$c \in !X \multimap (!X \otimes !X) \qquad w \in !X \multimap 1$$

corresponding to contraction and weakening. We can see these two maps as defining a structure of *comonoid* : intuitively this means the contraction map behaves like a commutative/associative law and that the weakening map behaves like its neutral element. The only difference with a usual monoid is that the arrows are in the wrong direction. A comonoid is therefore a triple (X, c, w) satisfying conditions of (co)-associativity, commutativity and neutrality. There are many examples of monoids among coherent spaces, since monoids are closed under $\otimes, \oplus$ and existential quantification (this means that given monoids, the above constructions can canonically be endowed with monoidal structures).

Let us call them *positive correlation spaces*.

Dually, spaces $?X$ are equipped with maps :

A	B	$A \wedge B$	$A \vee B$	$A \Rightarrow B$	$\perp A$	$\forall x A$	$\exists x A$
+1	+1	+1	+1	-1	-1	-1	+1
-1	+1	+1	-1	+1	+1	-1	+1
+1	-1	+1	-1	-1			
-1	-1	-1	-1	-1			

Table 1: Polarities for classical connectives.

$$c \in (?X \wp ?X) \multimap ?X$$

$$w \in \perp \multimap ?X$$

enjoying dual conditions, and that should be called “cocomonoids”, but we prefer to call them *negative correlation spaces*¹⁰. Negative correlation spaces are closed under $\wp, \&$ and universal quantification.

The basic idea to interpret classical logic will be to assign *polarities* to formulas, positive or negative, so that a given formula will be interpreted by a correlation space of the same polarity. The basic idea behind this assignment is that a positive formula has the right to structural rules on the left and a negative formula has the right to structural rules on the right of sequents. In other terms, putting everything to the right, either A or $A^\perp$ has structural rules for free. A classical sequent $\vdash \Gamma, \Delta$ with the formulas in Γ positive and the formulas in Δ negative is interpreted in linear logic as $\vdash ?\Gamma, \Delta$: the symbol $?$ in front of Γ is here to compensate the want of structural rule for positive formulas.

This induces a denotational semantics for classical logic. However, we easily see that there are many choices (using the two conjunctions and the two exponentials) when we want to interpret classical conjunction, similarly for disjunction, see [8], this volume. However, we can restrict our attention to choices enjoying an optimal amount of denotational isomorphisms. This is the reason behind the tables shown on next page.

It is easily seen that in terms of isomorphisms, negation is involutive, conjunction is commutative and associative, with a neutral element $\mathbf{V}$ of polarity +1, symmetrically for disjunction. Certain denotational distributivities $\wedge/\vee$ or $\vee/\wedge$ are satisfied, depending on the respective polarities.

Polarities are obviously a way to cope with the basic undeterminism of classical logic, since they operate a choice between the basic protocols of cut-

10. The dual of a comonoid is not a monoid

A	B	$A \wedge B$	$A \vee B$	$A \Rightarrow B$	$\perp A$	$\forall x A$	$\exists x A$
$+1$	$+1$	$A \otimes B$	$A \oplus B$	$A \multimap ?B$	$A^\perp$	$\forall x ?A$	$\exists x A$
-1	$+1$	$!A \otimes B$	$A \wp ?B$	$A^\perp \oplus B$	$A^\perp$	$\forall x A$	$\exists x !A$
$+1$	-1	$A \otimes !B$	$?A \wp B$	$A \multimap B$			
-1	-1	$A \& B$	$A \wp B$	$!A \multimap B$			

Table 2: Classical connectives : definition in terms of linear logic.

elimination. However, this is still not enough to provide a deterministic version of Gentzen's classical calculus **LK**. The reason lies in the fact that the rule of introduction of conjunction is problematic : from cliques in respectively $?X$ and $?Y$, when both X and Y are positive, there are two ways to get a clique in $?(X \otimes Y)$. This is why one must replace **LK** with another calculus **LC**, see [13] for more details, in which a specific positive formula may be distinguished. **LC** has a denotational semantics, but the translation from **LK** to **LC** is far from being deterministic. This is why we consider that our approach is still not absolutely convincing ... typically one cannot exclude the existence of a non-deterministic denotational semantics for classical logic, but God knows how to get it !

LC is indeed fully compatible with linear logic : it is enough to add a new polarity 0 (neutral) for those formulas which are not canonically equipped with a structure of correlation space. The miracle is that this combination of classical with intuitionistic features accommodates intuitionistic logic for free, and this eventually leads to the system **LU** of *unified logic*, see [14].

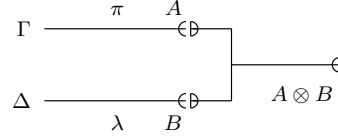
2.3 Geometry of interaction

At some moment we indicated an electronic analogy ; in fact the analogy was good enough to explain step (1) of cut-elimination (see subsection 1.3.6 by the fact that an extension cord has no action (except perhaps a short delay, which corresponds to the cut-elimination step). But what about the other links ?

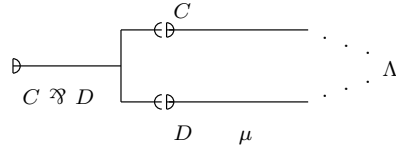
Let us first precise the nature of our (imaginary) plugs ; there are usually several pins in a plug. We shall restrict ourselves to one-pin plugs ; this does not contradict the fact that there may be a huge variety of plugs, and that the only allowed plugging is between complementary ones, labelled A and $A^\perp$.

The interpretation of the rules for $\otimes$ and $\wp$ both use the following well-known fact : two pins can be reduced to one (typical example : stereophonic broadcast).

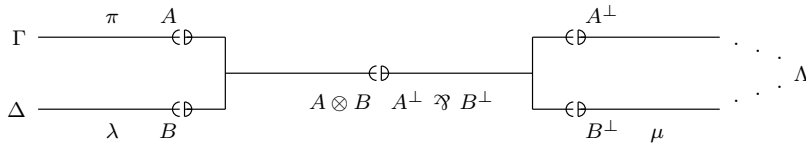
- $\otimes$ -rule : from units π, λ with respective interfaces $\vdash \Gamma, A$ and $\vdash \Delta, B$, we can built a new one by merging plugs A and B into another one (labelled $A \otimes B$) by means of an encoder.



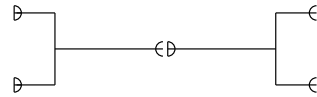
- $\wp$ -rule : from a unit μ with an interface $\vdash C, D, \Lambda$, we can built a new one by merging plugs C and D into a new one (labelled $C \wp D$) by means of an encoder :



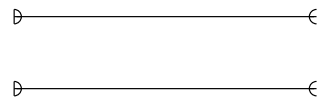
To understand what happens, let us assume that $C = A^\perp$, $D = B^\perp$; then $A^\perp \wp B^\perp = (A \otimes B)^\perp$, so there is the possibility of plugging. We therefore obtain



But the configuration



is equivalent to (if the coders are the same)



and therefore our plugging can be mimicked by two pluggings

$$\begin{array}{c}
 \Gamma \xrightarrow[\pi]{A \quad A^\perp} \epsilon \mathfrak{D} \xrightarrow{\quad} \dots \xrightarrow{\quad} \Lambda \\
 \Delta \xrightarrow[\lambda]{\quad} \epsilon \mathfrak{D} \xrightarrow[\mu]{B \quad B^\perp} \dots \xrightarrow{\quad} \Lambda
 \end{array}$$

If we interpret the encoder as $\otimes$ - or $\wp$ -link, according to the case, we get a very precise modelization of cut-elimination in proof-nets. Moreover, if we remember that coding is based on the development by means of Fourier series (which involves the Hilbert space) everything that was done can be formulated in terms of operator algebras. In fact the operator algebra semantics enables us to go beyond multiplicatives and quantifiers, since the interpretation also works for exponentials. We shall not go into this, which requires at least some elementary background in functional analysis ; however, we can hardly resist mentioning the formula for cut-elimination

$$\text{EX}(u, \sigma) := (1 - \sigma^2)u(1 - \sigma u)^{-1}(1 - \sigma^2)$$

which gives the interpretation of the elimination of cuts (represented by σ) in a proof represented by u . Termination of the process is interpreted as the nilpotency of σu , and the part $u(1 - \sigma u)^{-1}$ is a candidate for the execution. See [15], this volume, for more details. One of the main novelties of this paper is the use of *dialects*, i.e. data which are defined up to isomorphism. The distinction between the two conjunctions can be explained by the possible ways of merging dialects : this is a new insight in the theory of parallel computation.

Geometry of interaction also works for various λ -calculi, for instance for pure λ -calculus, see [7, 26]. It has also been applied to the problem of optimal reduction in λ -calculus, see [20].

Let us end this chapter by yet another refutation of weakening and contraction :

1. If we have a unit with interface $\vdash \Gamma$, it would be wrong to add another plug A ; such a plug (since we know nothing about the inside of the unit) must be a mock plug, with no actual connection with the unit ... Imagine a plug on which it is written “danger, 220V”, you expect to get some result if you plug something with it : here nothing will happen !
2. If we have a unit with a repetitive interface $\vdash \Gamma, A, A$, it would be wrong to merge the two similar plugs into a single one : in real life, we have such a situation with the stereophonic output plugs of an amplifier, which have exactly the same specification. There is no way to merge these two plugs

into one and still respect the specification. More precisely, one can try to plug a single loudspeaker to the two outputs plugs simultaneously ; maybe it works, maybe it explodes, but anyway the behaviour of such an experimental plugging is not covered by the guarantee ...

2.4 Game semantics

Recently Blass introduced a semantics of linear logic, see [5], this volume. The semantics is far from being complete (i.e. it accepts additional principles), but this direction is promising.

Let us forget the state of the art and let us focus on what could be the general pattern of a convincing game semantics.

2.4.1 Plays, strategies etc.

Imagine a game between two players **I** and **II** ; the rule determines which is playing first, and it may happen that the same player plays several consecutive moves. The game eventually terminates and produces a numerical output for both players, e.g. a real number. There are some distinguished outputs for **I** for which he is declared to be the winner, similarly for **II**, but they cannot win the same play. Let us use the letter σ for a strategy for player **I**, and the letter τ for a strategy for **II**. We can therefore denote by $\sigma * \tau$ the resulting play and by $< \sigma, \tau >$ the output. The idea is to interpret formulas by games (i.e. by the rule), and a proof by a winning strategy. Typically linear negation is nothing but the interchange of players etc.

2.4.2 The three layers

We can consider three kinds of invariants :

1. Given the game A , consider all inputs for **I** of all possible plays : this vaguely looks like a phase semantics (but the analogy is still rather vague) ;
2. Given the game A and a strategy σ for **I** consider the set $|\sigma|$ of all plays $\sigma * \tau$, when τ varies among all possible strategies for **II**. This is an analogue of denotational semantics : we could similarly define the interpretation $|\tau|$ of a strategy for **II** and observe that $|\sigma| \cap |\tau| = \{\sigma * \tau\}$ (this is analogue to the fact that a clique in X and a clique in $X^\perp$ intersect in at most one point) ;
3. We could concentrate on strategies and see how they dynamically combine : this is
analogous to geometry of interaction.

Up to the moment this is pure science-fiction. By the way we are convinced that although games are a very natural approach to semantics, they are not

primitive, i.e. that the game is rather a phenomenon, and that the actual semantics is a more standard mathematical object (but less friendly). Anyway, whatever is the ultimate status of games w.r.t. logic, this is an essential intuition : typically game semantics of linear logic is the main ingredient in the recent solution of the problem of full abstraction for the language **PCF**, see [1].

2.4.3 The completeness problem

The main theoretical problem at stake is to find a complete semantics for (first order) linear logic. Up to now, the only completeness is achieved at the level of provability (by phase spaces) which is rather marginal. Typically a complete game semantics would yield winning strategies only for those formulas which are provable. The difficulty is to find some semantics which is not contrived (in the same way that the phase semantics is not contrived : it does uses, under disguise, the principles of linear logic).

A non-contrived semantics for linear logic would definitely settle certain general questions, in particular which are the possible rules. It is not to be excluded that the semantics suggests tiny modifications of linear rules (e.g. many semantics accept the extra principle

$A \otimes B \multimap A \wp B$, known as *mix*), (and which can be written as a structural rule), or accepts a wider spectrum of logics (typically it could naturally be non-commutative, and then set up the delicate question of non-commutativity in logic). Surely it would give a stable foundation for constructivity.

Acknowledgements

The author is deeply indebted to Daniel Dzierzgowski and Philippe de Groote for producing L^AT_EX versions of a substantial part of this text, especially figures. Many thanks also to Yves Lafont for checking the final version.

BIBLIOGRAPHY

- [1] S. Abramsky, R. Jagadeesan, and P. Malacaria. Full abstraction for pcf (extended abstract). In Masami Hagiya and John C. Mitchell, editors, *Theoretical Aspects of Computer Software. International Symposium, TACS 94, Sendai, Japan*. Springer Verlag, April 1994. Lecture Note in Computer Science 789.
- [2] V. M. Abrusci. Non-commutative proof-nets. In *this volume*, 1995.
- [3] J.-M. Andreoli and R. Pareschi. Linear objects: logical processes with built-in inheritance. *New Generation Computing*, 9(3-4):445-473, 1991.
- [4] A. Asperti. A logic for concurrency. Technical report, Dipartimento di Informatica, Pisa, 1987.
- [5] A. Blass. A game semantics for linear logic. *Annals of Pure and Applied Logic*, 56:183-220, 1992.
- [6] R. Blute. Linear logic, coherence and dinaturality. *Theoretical Computer Science*, 93:3-41, 1993.
- [7] V. Danos. *La logique linéaire appliquée à l'étude de divers processus de normalisation et principalement du λ -calcul*. PhD thesis, Université Paris VII, 1990.
- [8] V. Danos, J.-B. Joinet, and H. Schellinx. LKQ and LKT : sequent calculi for second order logic based upon dual linear decompositions of classical implication. In *this volume*, 1995.
- [9] V. Danos and L. Regnier. The structure of multiplicatives. *Archive for Mathematical Logic*, 28:181-203, 1989.
- [10] T. Ehrhard. Hypercoherences : a strongly stable model of linear logic. In *this volume*, 1995.
- [11] C. Fouqueré and J. Vauzeilles. Inheritance with exceptions : an attempt at formalization with linear connectives in unified logic. In *this volume*, 1995.
- [12] J.-Y. Girard. Linear logic. *Theoretical Computer Science*, 50:1-102, 1987.
- [13] J.-Y. Girard. A new constructive logic : classical logic. *Mathematical structures in Computer Science*, 1:255-296, 1991.
- [14] J.-Y. Girard. On the unity of logic. *Annals of Pure and Applied Logic*, 59:201-217, 1993.

- [15] J.-Y. Girard. Geometry of interaction III : accommodating the additives. In *this volume*, 1995.
- [16] J.-Y. Girard. Light linear logic. *in preparation*, 1995.
- [17] J.-Y. Girard. Proof-nets : the parallel syntax for proof-theory. In *Logic and Algebra*, New York, 1995. Marcel Dekker.
- [18] J.-Y. Girard, Y. Lafont, and P. Taylor. *Proofs and types*, volume 7 of *Cambridge tracts in theoretical computer science*. Cambridge University Press, 1990.
- [19] J.-Y. Girard, A. Scedrov, and P.J. Scott. Bounded linear logic: A modular approach to polynomial time computability. *Theoretical Computer Science*, 97:1–66, 1992.
- [20] G. Gonthier, M. Abadi, and J.-J. Levy. The geometry of optimal lambda-reduction. In ACM Press, editor, *POPL'92*, Boston, 1992. Birkhäuser.
- [21] Y. Lafont. The linear abstract machine. *Theoretical Computer Science*, 59:95–108, 1990.
- [22] Y. Lafont. From proof-nets to interaction nets. In *this volume*, 1995.
- [23] J. Lambek. Bilinear logic in algebra and linguistics. In *this volume*, 1995.
- [24] P. Lincoln. Deciding provability of linear logic fragments. In *this volume*, 1995.
- [25] P. Lincoln, J. Mitchell, A. Scedrov, and N. Shankar. Decision problems for propositional linear logic. *Annals of Pure and Applied Logic*, 56:239–311, 1992.
- [26] Laurent Regnier. *Lambda-Calcul et Réseaux*. Thèse de doctorat, Université Paris 7, 1992.
- [27] D.N. Yetter. Quantales and non-commutative linear logic. *Journal of Symbolic Logic*, 55:41–64, 1990.