

PROOF-NETS : THE PARALLEL SYNTAX FOR PROOF-THEORY

Jean-Yves Girard

Institut de Mathématiques de Luminy, UPR 9016 – **CNRS**
163, Avenue de Luminy, Case 930, F-13288 Marseille Cedex 09

girard@iml.univ-mrs.fr

Abstract

The paper is mainly concerned with the extension of proof-nets to additives, for which the best known solution is presented. It proposes two cut-elimination procedures, the *lazy* one being in linear time. The solution is shown to be compatible with quantifiers, and the structural rules of exponentials are also accommodated.

Traditional proof-theory deals with cut-elimination ; these results are usually obtained by means of sequent calculi, with the consequence that 75% of a cut-elimination proof is devoted to endless commutations of rules. It is hard to be happy with this, mainly because :

- ▶ the structure of the proof is blurred by all these cases ;
- ▶ whole forests have been destroyed in order to print the same routine lemmas ;
- ▶ this is not extremely elegant.

However old-fashioned proof-theory, which is concerned with the ritual question : “is-that-theory-consistent ?” never really cared. The situation changed when subtle algorithmic aspects of cut-elimination became prominent : typically the determinism of cut-elimination, its actual complexity, its implementation cannot be handled in terms of sequent calculus without paying a heavy price. Natural deduction could easily fix the main drawbacks of cut-elimination, but this improvement was limited to the negative fragment of intuitionistic logic.

The situation changed in 1986 with the invention of linear logic : *proof-nets* were introduced in [G86] as a new kind of syntax for linear logic, in order to

cope with the problems arising from the intrinsic parallelism of linear sequent calculus. 9 years later the technology is perfectly efficient and proof-nets are now available for full linear logic¹. Using implicit translations, proof-nets are also available for classical and intuitionistic logics, i.e. for extant logical systems.

The essential result of [G86] was a *sequentialization* theorem showing the equivalence of the new syntax with the traditional one. However this original result was restricted to the multiplicative fragment **MLL**, non-multiplicative features being accommodated by *boxes*, i.e. sequents under disguise. Later on, the original sequentialisation proof was extended to quantifiers in [G88, G90] ; the second version, the only satisfactory one, made a significant use of simplifications of the method of [G86] discovered by Danos & Regnier, [DR88]. The problem of the extension to full linear logic, and especially to the additives (let us say to the fragment **MALL** of multiplicative-additive linear logic) has remained open for years. In fact two distinct problems must be solved :

- *To find the right notion of proof-net* : the notion was found quickly -say in the beginning of 1987- and since that time never really changed. The problem is to cope with the $\&$ -rule of sequent calculus, for which a superimposition of two proof-nets must be made ; by introducing for each $\&$ -link a boolean variable which distinguishes between the two *slices* of the superimposition, one eventually gets a notion of net in which formulas and links are *weighted* by boolean polynomials. Although this notion remains the only serious candidate, it is to be noted that it is far from being absolutely satisfactory : this is because the question of determining the identity between two formulas (or two links) in two different proof-nets cannot receive a satisfactory answer, especially out of **MALL**.
- *To prove sequentialisation* : the case of quantifiers was considered as a preliminary case, but the solution found in [G88] was too acrobatic to be extended. This solution was improved in [G90] by means of a new kind of switchings, determined by certain *formal* dependencies. The solution immediately extends to additives, provided certain *weight* dependencies are forbidden (this is the *dependency condition*, which requires all weights to be monomials). Since 1990 we have been trying to get rid of this technical limitation...and we must admit that this was a stupid attitude : in spite of their limitations our proof-nets obtained were extremely efficient (elimination of boxes is always a big simplification), whereas getting rid of the dependency restriction would of course make the proof-nets more intrinsic in some cases, but would still not make them absolutely satisfactory. This is the reason why we eventually decided to publish our partial results.

1. Except the additive neutrals

We shall first develop the present version of additive proof-nets, which is enough for most applications, together with a restricted (*lazy*) cut-elimination. In an appendix we shall consider possible improvements and the general cut-elimination procedure.

Of course, proof-nets are not intended for multiplicative and additives only. In fact we shall devote some sections to the (unproblematic) extension to quantifiers, and to the structural rules in the exponential case (here with limited problems with weakening). Therefore we get rid of all boxes, but exponential boxes. This is of course a tremendous improvement over traditional sequent calculus.

1 PROOF-NETS FOR MALL

1.1 Proof-structures

Definition 1

A link L is an expression

$$\frac{P_1, \dots, P_n}{Q_1, \dots, Q_m} L$$

involving n formulas (the premises of L) $P_1, \dots, P_n$ and m formulas (the conclusions of L) $Q_1, \dots, Q_m$

ID –links	: 0 premise	2 conclusions : $A, A^\perp$
CUT –links	: 2 premises : $A, A^\perp$	0 conclusion
$\otimes$ –links	: 2 premises : A, B	1 conclusion : $A \otimes B$
$\wp$ –links	: 2 premises : A, B	1 conclusion : $A \wp B$
$\oplus_1$ –links	: 1 premise : A	1 conclusion : $A \oplus B$
$\oplus_2$ –links	: 1 premise : B	1 conclusion : $A \oplus B$
$\&$ –links	: 2 premises : A, B	1 conclusion : $A \& B$

The premises of $\otimes, \wp, \&$ -links are ordered : this means that we can distinguish a left premise (here A) and a right premise (here B). On the other hand the premises of a CUT -link and the conclusions of an ID -link are unordered.

Remark. — It is convenient to consider generalized axioms $\vdash A_1, \dots, A_n$ ($n > 0$), which are interpreted by generalized axiom links (no premise, but ordered conclusions $A_1, \dots, A_n$). Such generalized axioms will occur in the proof of our main theorem 2 ; they also occur when one wants to accommodate other styles of syntax, which are foreign to the proof-net technology, in which case they are called *boxes*. The idea of a box is that from the outside it looks like a generalized axiom, whereas it has an inside which can be in turn another proof-net. A box freezes n formulas,

and can therefore be seen as a sequent, the conclusion of a rule, whose premises are proven in the box. Traditional sequent calculus is therefore a system of proof-nets in which the only links are boxes, and all the improvement made in 9 years consist in progressively restricting the use of boxes : in this paper boxes are limited to exponential connectives (and to the neutral $\top$).

Remark. — One should never speak of formulas, but of *occurrences*, which is extremely awkward. We adopt once for all the convention that all our formulas are distinct (for instance by adding extra indices). In particular $ID, \otimes, \wp, \oplus, \&$ -links are determined by their conclusion(s), and a CUT -link is determined by its premises.

Definition 2

- If L is a $\&$ -link, with $A \& B$ as its conclusion, we introduce the eigenweight p_L , which is a boolean variable. The intuitive meaning of p_L is the choice l/r between the two premises A and B of the link, p_L for “left”, i.e. A , $\neg p_L$ for “right”, i.e. B ; we use ϵp_L to speak of p_L or $\neg p_L$.
- If Θ is a structure involving the $\&$ -links $L_1, \dots, L_k$ (with associated eigenweights $p_1, \dots, p_k$), then a weight (relative to Θ) is any element of the boolean algebra generated by $p_1, \dots, p_k$.

Definition 3

A proof-structure Θ consists of :

- A set of formulas (see the previous remark) ;
- A set of links ; each of these links takes its premise(s) and conclusion(s) among the formulas of Θ ;
- For each formulas A of Θ , a weight $w(A)$, i.e. a non-zero element of the boolean algebra generated by the eigenweights $p_1, \dots, p_n$ of the $\&$ -rules of Θ ;
- For each link L of Θ , a weight $w(L)$.

satisfying the following conditions :

- Each formula is the the premise of at most one link and the conclusion of at least one link ; the formulas which are not premises of some link are called the conclusions of Θ ;
- $w(A) = \sum w(L)$, the sum being taken over the set of links with conclusion A ;
- if A is a conclusion of Θ , then $w(A) = 1$;
- if w is any element of the boolean algebra generated by the weights occurring in Θ , and L is a $\&$ -link, then $w.\neg w(L)$ does not depend on p_L , i.e. belongs to the boolean algebra generated by the eigenweights distinct from p_L ;

- if w is any weight occurring in Θ , then w is a monomial $\epsilon_1 p_{L_1} \dots \epsilon_k p_{L_k}$ of eigenweights and negations of eigenweights² ;
- $w(L) \neq 0$; moreover if L is any non-identity link, with premises A and (or) B then
 - if L is any of $\otimes, \wp, CUT$, then $w(L) = w(A) = w(B)$;
 - if L is $\oplus_1$, then $w(L) = w(A)$;
 - if L is $\oplus_2$, then $w(L) = w(B)$;
 - if L is a $\&$ -link, then $w(A) = w(L).p_L$ and $w(B) = w(L).\neg p_L$, (hence $w(L) = w(A) + w(B)$).

Remark. —

- Weights are in a boolean algebra, and therefore both algebraic and logical graphism can be used ; here we decide to use the product notation (instead of the intersection), but we keep $\neg w$ (instead of $1 - w$) ; when we use the sum, we of course mean the disjoint union, i.e. when I write $w(L) = w(A) + w(B)$, I implicitly mean that $w(A).w(B) = 0$.
- The technical condition “ $w.\neg w(L)$ does not depend on p_L ” says that the boolean variable p_L has no real meaning “outside $w(L)$ ” ; applying the condition to $\neg w(L)$, we see that $w(L)$ does not depend on p_L , in particular $w(L).\epsilon p_L \neq 0$.
- There are two ways to think of the dependency condition : either as a technical restriction needed for the sequentialisation theorem (all our efforts to get rid of it failed) or as a nice companion to the previous condition, since both are very natural when a proof-structure is seen as a coherent space, see A.1.1.

1.2 Sequent calculus and proof-nets

Definition 4

Let Θ be a proof-structure and let L be either a *CUT*-link, or a link with only one conclusion, which is in turn a conclusion of Θ and such that $w(L) = 1$; we say that L is a *terminal link* of Θ . Given such a link, we define the removal of L in Θ which consists (provided it makes sense) in one or two proof-structures.

- If L is a $\otimes$ -link (resp. a *CUT*-link) with premises A, B , and $\Gamma, A \otimes B$ (resp. Γ) is the set of conclusions of Θ : the removal of L consists in partitioning (if possible) the formulas of Θ distinct from $A \otimes B$ (resp. the formulas of Θ) in two subsets X and Y , one containing A , the other containing B , in such a way that, whenever a link L' distinct from L has a premise or a conclusion in X (resp. in Y), then all other premises

2. This is the *dependency* condition

and conclusions of L' belong to X (resp. to Y). The restrictions Θ/X and Θ/Y are defined in an obvious way, and are proof-structures with respective conclusions Γ', A and Γ'', B . Observe that $\Gamma = \Gamma', \Gamma''$.

- If L is a $\wp$ -link with premises A, B , and $\Gamma, A \wp B$ is the set of conclusions of Θ : the removal of L consists in removing the conclusion $A \wp B$ and the link L ; this induces a proof-structure with conclusions Γ, A, B .
- If L is a $\oplus_1$ -link with premise A , and $\Gamma, A \oplus B$ is the set of conclusions of Θ : the removal of L consists in removing the conclusion $A \oplus B$ and the link L ; this induces a proof-structure with conclusions Γ, A .
- If L is a $\oplus_2$ -link with premise B , and $\Gamma, A \oplus B$ is the set of conclusions of Θ : the removal of L consists in removing the conclusion $A \oplus B$ and the link L ; this induces a proof-structure with conclusions Γ, B .
- If L is a $\&$ -link with premises A, B , and $\Gamma, A \& B$ is the set of conclusions of Θ : the removal of L consists in first removing the conclusion $A \& B$ and the link L (to get Θ') and then forming two proof-structures Θ_A and Θ_B :
 - In Θ' make the replacement $p_L = 1$, and keep only those links L' whose weight is still non-zero, together with the premises and conclusions of such links : the result is by definition Θ_A , a proof-structure with conclusions Γ, A .
 - In Θ' make the replacement $p_L = 0$, and keep only those links L' whose weight is still non-zero, together with the premises and conclusions of such links : the result is by definition Θ_B , a proof-structure with conclusions Γ, B .

Definition 5

A proof-structure Θ is sequentialisable when it can be reduced, by iterated removal of terminal rules, to identity links. In more pedantic terms :

- An identity link is sequentialisable ;
- If the result of removing the terminal link L in Θ yields sequentialisable proof-structures, then Θ is sequentialisable.

Remark. —

- The removal of a given terminal link is not always possible, and its result is not necessarily unique (however, for *proof-nets*, it would be easy to show, by a connectivity argument, that the removal of a $\otimes$ - or *CUT*-link is unique).
- Each removal step consists in the writing down of a rule of MALL ; therefore a sequentialisable proof-structure has a *sequentialisation*, which consists in a proof in MALL.
- Conversely, given a proof Π of $\vdash \Gamma$ in sequent calculus, one can build a proof-structure Π° with conclusions Γ , such that Π is a sequentialisation of Π° . But contrarily to the situation of our previous papers [G86, G88, G90] on proof-nets, Π° is no longer unique. The problem is in the interpretation of a $\&$ -rule

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

applied to proofs Π_1 and Π_2 of $\vdash \Gamma, A$ and $\vdash \Gamma, B$. We must find a proof-structure Π° such that the removal of a terminal $\&$ -link yields Π_1° and Π_2° . The basic idea is to merge the two proof-structures by means of the eigenweight p_L : anything with respective weights w_1 and w_2 in Π_1° and Π_2° will now get the weight $p_L.w_1 + \neg p_L.w_2$ (we give the weight $w_1 = 0$ to something which is absent in Π_1°). This simple idea yields the following list of cases :

- If X is a formula or a link occurring in both of Π_1° and Π_2° , with the weights $w_1(X)$ and $w_2(X)$, then X will occur in Π° with the weight $p_L.w_1(X) + \neg p_L.w_2(X)$; in particular the formulas of Γ occur in Π° with the weight 1.
- If X is a formula or a link occurring in Π_1° but not in Π_2° , with the weight $w_1(X)$, then X will occur in Π° with the weight $p_L.w_1(X)$; in particular, A will occur with the weight p_L .
- If X is a formula or a link occurring in Π_2° but not in Π_1° , with the weight $w_2(X)$, then X will occur in Π° with the weight $\neg p_L.w_2(X)$; in particular, B will occur with the weight $\neg p_L$.

and to add to this merge the formula $A \& B$ together with the $\&$ -link L whose premises are A, B and whose conclusion is $A \& B$; both L and $A \& B$ receive the weight 1. But this is not as simple as it might seem : how do we know that a formula or a link X of Π_1 is the same as another formula or link Y of Π_2 ? There is no simple answer (except in some very specific cases, see discussion A.1.4), and moreover, in cases where we feel entitled to make such identifications, the resulting weight $p_L.w_1(X) + \neg p_L.w_2(X)$ is not a monomial, which contradicts the dependency condition. However there is at least the possibility to decide that no identification between Π_1 and Π_2 is made, but for the conclusions, i.e. the formulas of Γ ; by the way the sequent calculus formulation of the $\&$ -rule

stipulates that the contexts of the two premises must be equal, hence this is a clear case where there is no doubt as to the identification between a formula in Π_1 and a formula in Π_2 .

Anyway, the main problem is to find a *sequentialisation* theorem ; this means to give an intrinsic characterization of *sequentialisable* proof-structures. The answer (a notion of *proof-net*) to this problem is important because we want to carry our program of geometrization of proofs, introduced in [G87A] under the name of *geometry of interaction*. Remember that *geometry of interaction* is issued from the analysis of multiplicative proof-nets in terms of permutations, [G86A]. In fact the redaction of this paper was postponed until a geometry of interaction for additives was found, which is now the case, see [G94].

1.3 A wrong answer : slicing

Definition 6

Let φ be a valuation for Θ , i.e. a function from the set of eigenweights of Θ into the boolean algebra $\{0, 1\}$, which induces a function (still denoted φ) from the weights of Θ to $\{0, 1\}$. The slice $\varphi(\Theta)$ is obtained by restricting to those formulas A of Θ such that $\varphi(w(A)) = 1$, with an obvious modification for the remaining $\&$ -links : only one premise is present.

The definition suggests a simple-minded criterion, (which is anyway an approximation to the real solution) : observe that (if we neglect unary links), $\varphi(\Theta)$ is a multiplicative proof-structure. Therefore we can require $\varphi(\Theta)$ to be multiplicatively correct, i.e. to be a multiplicative proof-net. This condition is obviously necessary (in fact it is what we get by restricting definition 10 to *normal* jumps).

But the condition cannot be sufficient : it corresponds to a separate treatment of additives and multiplicatives, without real interaction between the conditions. This would mean that all multiplicatives distribute over all additives : for instance the obvious proof-structure with conclusions

$(A^\perp \wp B^\perp) \oplus (A^\perp \wp C^\perp)$ and $A \otimes (B \& C)$ has two multiplicatively correct slices, although it states a wrong instance of distributivity, and is not sequentialisable, whereas the obvious proof-structure with conclusions $(A^\perp \wp B^\perp) \& (A^\perp \wp C^\perp)$ and $A \otimes (B \oplus C)$ (which has essentially the same slices) is sequentialisable. Our ultimate criterion must therefore force the additive and multiplicative layers to interact.

1.4 Proof-nets

Our basic idea will be to mimic our criterion of [G90] ; in this paper, certain switchings for $\forall$ -links were induced by the dependency of some formula upon

the *eigenvariable* of the link.

Definition 7

Let φ be a valuation of Θ , let p_L be an eigenweight ; we say that the weight w (in Θ) depends on p_L (in $\varphi(\Theta)$) iff $\varphi(w) \neq \varphi_L(w)$, where the valuation φ_L is defined by :

- $\varphi_L(p_L) = \neg(\varphi(p_L))$
- $\varphi_L(p_{L'}) = \varphi(p_{L'})$ if $L' \neq L$.

A formula A of Θ is said to depend on p_L (in $\varphi(\Theta)$), if A is conclusion of a link L' such that $\varphi(w(L')) = 1$ and $\varphi_L(w(L')) = 0$. This basically means that A and L' are present in $\varphi(\Theta)$, but that changing the value of the valuation for p_L would make A (or at least L') disappear from the slice.

Definition 8

A switching $\mathcal{S}$ of a proof-structure Θ consists in :

- The choice of a valuation $\varphi_{\mathcal{S}}$ for Θ ;
- The selection of a choice $\mathcal{S}(L) \in l, r$ for all $\wp$ -links of $\varphi_{\mathcal{S}}(\Theta)$;
- The selection for each $\&$ -link L of $\varphi_{\mathcal{S}}(\Theta)$ a formula $\mathcal{S}(L)$, the jump of L , depending on p_L in $\varphi_{\mathcal{S}}(\Theta)$. There is always a normal choice of jump for L , namely the premise A of L such that $\varphi_{\mathcal{S}}(w(A)) = 1$.

Definition 9

Let $\mathcal{S}$ be a switching of a proof-structure Θ ; we define the graph $\Theta_{\mathcal{S}}$ as follows :

- The vertices of $\Theta_{\mathcal{S}}$ are the formulas of $\varphi_{\mathcal{S}}(\Theta)$;
- For all *ID*-links of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the conclusions ;
- For all generalized axiom links with conclusions $A_1, \dots, A_n$, we draw an edge between A_1 and A_2 , etc., A_{n-1} and A_n ;
- For all *CUT*-links of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the premises ;
- For all $\oplus$ -links of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the conclusion and the premise ;
- For all $\otimes$ -links of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the left premise and the conclusion, and between the right premise and the conclusion ;
- For all $\wp$ -links L of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the premise (left or right) selected by $\mathcal{S}(L)$ and the conclusion ;
- For all $\&$ -links L of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the jump $\mathcal{S}(L)$ of L and the conclusion.

Definition 10

A proof-structure Θ is said to be a proof-net when for all switchings $\mathcal{S}$, the graph $\Theta_{\mathcal{S}}$ is connected and acyclic.

We immediately get the

Theorem 1

If Θ is sequentialisable, then Θ is a proof-net.

PROOF. — The proof is straightforward and uninteresting. $\square$

1.5 The sequentialisation theorem**Theorem 2**

Proof-nets are sequentialisable.

We shall devote the remainder of the subsection to the proof of the theorem. We shall make some simplifying hypotheses :

1.5.1 Empires

In this subsection and in the next one, we fix one for all a valuation φ and we concentrate on the slice $\varphi(\Theta)$, where Θ is a given proof-net. All switchings $\mathcal{S}$ considered are such that $\varphi_{\mathcal{S}} = \varphi$.

Definition 11

If $\mathcal{S}$ is a switching and A is a formula of $\varphi(\Theta)$, then we modify the graph $\Theta_{\mathcal{S}}$ by deleting (in case A is premise of a link L) any edge induced by L and connecting A to another formula B . (There is at most such an edge, connecting A with the conclusion of L or with the other premise of L if L is a CUT-link.) This determines a partition of $\Theta_{\mathcal{S}}$ into at most two connected components, and the one containing A is denoted $\Theta_{\mathcal{S}}^A$. We define the empire of A as $eA := \bigcap_{\mathcal{S}} \Theta_{\mathcal{S}}^A$, the intersection being taken over all switchings $\mathcal{S}$.

Lemma 1

Empires are imperialistic ; this means that, as soon B_1 and B_2 are linked by L and $B_1 \in eA$, then $B_2 \in eA$, with only three exceptions :

- B_1 is A and is a premise of L : if B_2 is the conclusion of L , it is not in eA ;
- L is a $\mathfrak{A}$ -link, B_2 is the conclusion of the link, B_1 is one of the two premises of L and the other premise of L is not in eA : in such a case the conclusion is never in eA ;

- L is a $\&$ -link, B_2 is the conclusion of the link, B_1 is the premise of L (remaining in $\varphi(\Theta)$) and there is a formula C in $\varphi(\Theta)$ depending on p_L which is not in eA : in such a case the conclusion is never in eA .

PROOF. — straightforward ; see e.g. [G90], lemma 1. We sketch the proof for the sake of self-containment :

- one first remarks that if B_1 is a conclusion of L , then $B_2 \in eA$; the typical example is that of a $\mathfrak{A}$ -link. If $\mathcal{S}$ is a switching inducing two components, and such that $B_2 \notin \Theta_{\mathcal{S}}^A$, then we see (perhaps by changing our mind about $\mathcal{S}(L)$) that $B_1 \notin \Theta_{\mathcal{S}}^A$. The other cases are handled in the same way. Observe that the same argument shows that if the conclusion B_1 of a $\&$ -link L is in eA , so is any formula whose weight depends on p_L .
- it only remains to show that if L is a $\mathfrak{A}$ - or $\&$ -link and all possible choices for $\mathcal{S}(L)$ are in eA , so is the conclusion B of L , provided $B \neq A$. But any switching $\mathcal{S}$ must connect B with one of these formulas which are in $\Theta_{\mathcal{S}}^A$, and this forces $B \in \Theta_{\mathcal{S}}^A$. $\square$

Lemma 2

There is an $\mathcal{S}$ such that $eA = \Theta_{\mathcal{S}}^A$; $\mathcal{S}$ is called a *principal switching* for eA .

PROOF. — $\mathcal{S}$ is obtained as follows : in case A is the premise of a $\mathfrak{A}$ - or $\&$ -link, set $\mathcal{S}$ so as to draw an edge between A and the conclusion of the link ; the other $\mathfrak{A}$ and $\&$ -links are switched as follows :

- if L is a $\mathfrak{A}$ -link with its conclusion C not in eA , then one of its premises, let us say B is not in eA by lemma 1 and we set $\mathcal{S}$ so as to draw an edge between B and C ;
- if L is a $\&$ -link L with its conclusion C not in eA , then a formula of $\varphi(\Theta)$, let us say B , depending on p_L , is not in eA by lemma 1 and we set $\mathcal{S}$ so as to draw an edge between B and C .

It is immediate that $eA = \Theta_{\mathcal{S}}^A$. $\square$

Lemma 3

Let A and B be distinct formulas in $\varphi(\Theta)$ and assume that $B \notin eA$; then

- if $A \in eB$, then $eA \subset eB$;
- if $A \notin eB$, then $eA \cap eB = \emptyset$.

PROOF. — see [G90], lemma 3. We also sketch the proof : we define a principal switching of eB as in lemma 2 by adding another constraint : if L is a $\mathfrak{A}$ or $\&$ whose conclusion is in eB , then in case A is a premise of L set $\mathcal{S}$ so as to connect the conclusion with A , and (otherwise) if some possible choice for $\mathcal{S}(L)$ is not in eA , make this very choice. Now if a formula C belongs to $eA \cap eB$, then C is connected with B (which is not in eA) inside $\Theta_{\mathcal{S}}^B = eB$, which means that some edge induced by $\mathcal{S}$ inside eB links a formula in eA with a formula outside eA . This is only possible (because of lemma 1) when the formula of eA is A . This proves the second part of the claim. If $A \in eB$, A must be the premise of a link L whose conclusion is not in eA , and $\mathcal{S}(L)$ has been set so as to get two connected components : $\Theta_{\mathcal{S}}^A$ and its complement. It is immediate that $\Theta_{\mathcal{S}}^A \subset \Theta_{\mathcal{S}}^B$ hence $eA \subset \Theta_{\mathcal{S}}^A \subset \Theta_{\mathcal{S}}^B = eB$. $\square$

Definition 12

A formula B of eA is said to be a door of eA iff

- *Either it is the premise of a link L whose conclusion does not belong to eA ;*
- *Or it is a conclusion of Θ . Obviously A is a door of eA , the main door ; the other doors are called auxiliary doors. The set of doors of eA is the border of eA .*

Lemma 4

Let C be an auxiliary door of eA which is not a conclusion of Θ ; then C is the premise of a $\mathfrak{A}$ - or a $\&$ -rule.

PROOF. — immediate from lemma 1. $\square$

1.5.2 Maximal empires

The hypotheses are the same as in the previous section ; but we now assume that A is the conclusion of a $\&$ -link L and that eA is maximal w.r.t. inclusion among similar empires (i.e. that if B is the conclusion of a $\&$ -link and $eA \subset eB$, then $A = B$).

Lemma 5

Let L' be any $\&$ -link and let $B, B' \in \varphi(\Theta)$ be such that B, B' depends on $p_{L'}$, and assume that $B \in eA$; then $B' \in eA$.

PROOF. — let w be the weight of the link whose conclusion is B and such that $\varphi(w) = 1$; then the technical condition “ $w.\neg w(L')$ does not depend on $p_{L'}$ ” ensures that $\varphi(\neg w(L')) = 0$, hence if C is the conclusion of L' , that $C \in \varphi(\Theta)$.

By lemma 1 $B \in eC$, hence $eA \cap eC \neq \emptyset$; by maximality $eA \subset eC$ is impossible, hence by lemma 3 we get $C \in eA$. Now by lemma 1 this implies $B' \in eA$. $\square$

Lemma 6

If C is a border formula of eA and L' is any $\&$ -link, then $\varphi_{L'}(w(C)) = 1$,
i.e. the border formulas are still present in $\varphi'_L(\Theta)$.

PROOF. — C is either a conclusion, in which case $w(C) = 1$ or C is a premise of some rule L'' : let B be the conclusion of L'' (in case L'' is a cut, let B be the other premise of L''). Then $B \notin eA$ but since B depends on L' , lemma 5 yields $B \in eA$, a contradiction. $\square$

1.5.3 Stability of maximal empires

The purpose of this section is to show that if A is the conclusion of a $\&$ -link and is maximal among similar empires w.r.t. a given φ , then it remains maximal w.r.t. any other choice of φ . It will be enough to start with a given φ and to show that A is still maximal w.r.t. φ_L . We can only prove this essential fact under the dependency condition, which has the following consequence :

Lemma 7

Assume that $\varphi(w(A)) = 1$ and that (w.r.t. φ) $w(A)$ depends neither on p_L and p_M ; then this remains true w.r.t. φ_L .

PROOF. — the monomial $w(A)$ cannot make use of ϵp_L etc. $\square$

This is wrong for non monomials, typically $p \cup q$ depends neither on p nor on q if $\varphi(p) = \varphi(q) = 1$, but depends on q if $\varphi(p) = 0$. This phenomenon is exactly the familiar failure of stability in the case of the *parallel or*.

Lemma 8

Assume that $B \in eA$ (w.r.t. φ) and that $w(B)$ does not depend on p_L ; then w.r.t. φ_L we still have $B \in eA$.

PROOF. — it takes nothing to assume that A is not a conclusion (in which case eA is $\varphi(\Theta)$ and everything is trivial) ; so one can fix a formula C with the same weight as A such that C cannot belong to eA (take C to be the conclusion of the link of which A is a premise, or the other premise if this link is L a *CUT*). Let D be the conclusion of the link L , then :

- either $D \in eA$, hence by lemma 5, the formulas outside eA do not depend on p_L , hence by lemma 7, changing φ to φ_L does not alter any dependency outside. Assume that a switching $\mathcal{S}$ has been chosen, (w.r.t. φ_L) with $B \in \Theta_{\mathcal{S}}^A$,

$C \notin \Theta_S^A$. Then if we define a switching $\mathcal{S}'$ (w.r.t. φ) essentially by making the same choices outside eA (which is possible, since the dependencies are unaffected), then B and C are still connected, hence we still have $B \in \Theta_{\mathcal{S}'}^A$, contradicting the hypothesis.

- or $D \notin eA$, hence by lemma 5, the formulas inside eA do not depend on p_L , hence by lemma 7, changing φ to φ_L does not alter any dependency inside. Assume that a switching $\mathcal{S}$ has been chosen, (w.r.t. φ_L) with $B \notin \Theta_S^A$. Then if we define a switching $\mathcal{S}'$ (w.r.t. φ_L) essentially by making the same choices inside eA (which is possible, since the dependencies are unaffected), then B and A are still not connected, hence we still have $B \notin \Theta_{\mathcal{S}'}^A$, contradicting the hypothesis. $\square$

Lemma 9

eA is maximal w.r.t. φ_L .

PROOF. — assume that (w.r.t. φ_L) $eA \subset eB$, with eB maximal. Then (still w.r.t. φ_L) $A \in eB$ but $B \notin eA$. Using the maximality of eA w.r.t. φ , of eB w.r.t. φ_L and lemma 8, we get (w.r.t. φ) $A \in eB$ and $B \notin eA$, which contradicts the maximality of eA w.r.t. φ . $\square$

1.5.4 Proof of the theorem

PROOF. — by induction on the number n of $\&$ -links in Θ :

- if $n = 0$, then we are basically in the multiplicative case and we are done. To be precise, we argue by induction on the number m of links of the proof-net :
 - if $m = 1$, then the proof-net consists in a single axiom link and we are done ;
 - if $m > 1$, then by connectedness Θ must contain some link which is not an axiom, hence some terminal link (see definition 4) exists, and
 - * if there is a terminal $\wp$ or $\oplus_i$ -link, then we can remove it, thus getting another proof-net (immediate) with a smaller value of m , to which the induction hypothesis applies ; then Θ is sequentialisable ;
 - * if all terminal links are $\otimes$ - or CUT -links, with premises A_i, B_i , choose one of these premises (say A_1) such that eA_1 is maximal inside the eA_i and eB_i . If C is a border formula of eA_1 and C is not a conclusion of Θ , then C stands (hereditarily) above some A_i or B_j , let us say above B_2 . Then by lemma 1, $C \in eB_2$, hence by lemma 3 $eA_1 \subset eB_2$, contradicting the maximality of eA_1 . This shows that the border of eA_1 only consists of conclusions. But then it is easy to see that $\Theta_S^{A_1}$ is always equal to

eA_1 (apply lemma 1), hence $\Theta_S^{B_1}$ is always equal to eB_1 . This means that the removal of our terminal link induces a splitting of the proof-net into two connected components, which are obviously proof-nets and to which the induction hypothesis applies ; Θ is therefore sequentialisable.

- if $n > 0$, then choose φ (e.g. $\varphi(p_L) = 1$ for all L), and let A be a conclusion of some $\&$ -link such that eA is maximal w.r.t. φ . Now by lemma 9 eA remains maximal w.r.t. to any φ' , and its border remains the same by lemma 6. Moreover, if L is any $\&$ link whose conclusion B belongs to eA w.r.t. φ , B still belongs to eA w.r.t. any φ' such that $\varphi'(w(L)) = 1$ (this is another use of the dependency condition : since $w(L)$ is a monomial, it is possible to write φ' as $\varphi_{L_1 \dots L_k}$ for a suitable sequence $L_1 \dots L_k$ with the property that all intermediate valuations $\varphi_{L_1 \dots L_i}$ yield the same value $\varphi_{L_1 \dots L_i}(w(L)) = 1$. Then lemma 8 is enough to conclude.) This means that a given eigenweight p_L cannot occur both in eA for some valuation and in the complement of eA for some other valuation. In other terms, we can split the eigenweights into two disjoint groups, those who may occur in eA (group I) and those who may occur in its complement (group II). Now let us introduce two new proof-nets :
 - a proof-net Θ' whose conclusions are those of the (constant) border of eA ; a formula B is in Θ' when $B \in eA$ for some valuation. This immediately induces a proof-structure (take those links whose conclusion is in Θ' , with the weight they had in Θ). Θ' contains all occurrences of eigenvariables from group I, and no occurrence from group II. It is immediate that Θ' is a proof-net. Moreover Θ' has a terminal $\&$ -link and removing this link as in definition 4 induces two proof-nets (immediate) to which the induction hypothesis applies, and which are therefore sequentialisable. Then Θ' is sequentialisable. In case A is a conclusion of Θ , then $\Theta = \Theta'$ and we are done ; otherwise we introduce
 - a proof-net Θ'' consisting of those formulas B such that $B \notin eA$ for some valuation φ' such that $\varphi'(w(B)) = 1$ and of the border formulas of eA which are not conclusions (there are some). The links of Θ'' are all the links of Θ whose conclusion is outside eA and a new axiom link whose conclusions consist in the border of eA . The links and formulas coming from Θ receive the same weight, whereas the new axiom link is weighted 1. Θ'' contains all occurrences of eigenvariables from group II, and no occurrence from group I. Now it is an easy exercise in graph theory to show that Θ'' is a proof-net. The induction hypothesis applies and Θ'' is therefore sequentialisable. By the way observe that if we replace our new axiom link in Θ'' by the proof-net Θ' , the result is the original Θ .

The sequentialisation of Θ is obtained by taking the sequentialisation of Θ'' and replacing the sequent calculus axiom corresponding to the new axiom link with the sequentialisation of Θ' . $\square$

1.6 Cut-elimination in MALL (lazy procedure)

There is a Church-Rosser cut-elimination procedure for additive proof-nets, which enjoys the subformula property. For practical uses this procedure is challenged by a coarser one, *lazy* cut-elimination, which can be proved to achieve the same result in some important cases. The more general procedure is described in appendix A.1.3.

1.6.1 Lazy cut-elimination

Lazy cut-elimination is concerned with only a restricted kind of cut :

Definition 13

A cut-link L is said to be ready iff

- *the weight of the cut-link is 1*
- *both premises of the cut are the conclusion of exactly one link (these links are therefore also of weight 1)*

Theorem 3

Let Θ be a proof-net whose conclusions do not contain the connective $\&$ and without ready cut ; then Θ is cut-free.

PROOF. — if Θ contains no $\&$ -link, we are done, since all weights are 1. Otherwise, choose a $\&$ -link L in such a way that the empire eA of its conclusion is maximal among similar empires (w.r.t. any valuation). Then A is weighted 1 as well his (hereditary) conclusions. Below A sits a terminal link L' with no conclusion (otherwise A would be a subformula of this conclusion, and $\&$ would occur in a conclusion), hence L' is a cut of weight 1. In fact A sits hereditarily above the premise B of L which is in turn the conclusion of a link of weight 1 (either L or a link "below" L). Now the premise $B^\perp$ of L' might be the conclusion of several links, in which case there is a valuation φ and a $\&$ -link L'' (with conclusion C) such that $B^\perp$ depends on $p_{L''}$. In the slice $\varphi(\Theta)$ observe that $B^\perp \notin eA$, but $A \in eC$ (this is because $B^\perp \in eC$ and so $B \in eC$ and then $A \in eC$ by lemma 1). But then $eA \subset eC$ by lemma 3, a contradiction. $\square$

Remark. —

- The theorem therefore establishes that the lazy procedure which only removes ready cuts is enough in the absence of additive connectives. In fact since the general procedure extends the lazy one and is Church-Rosser, the lazy procedure yields the same result in this important case.
- The theorem still holds for full linear logic : if a proof-net is without ready cuts and its conclusions mention neither $\&$ nor existential higher order quantifiers, then it is cut-free. The restriction on existentials is just to forbid the hiding of a $\&$ by a $\exists$ -rule, which can be the case with higher order.

Definition 14

Let L be a ready cut in a proof-net Θ , whose premises A and $A^\perp$ are the respective conclusions of links L, L' . Then we define the result Θ' of reducing our cut in Θ :

- If L is an ID -link then Θ' is obtained by removing in Θ the formulas A and $A^\perp$ (as well as the cut-link and L) and giving a new conclusion to L' : the other conclusion of L , which is another occurrence of $A^\perp$ (ID -reduction).
- If L is a $\otimes$ -link (with premises B and C) and L' is a $\wp$ -link (with premises $B^\perp$ and $C^\perp$), then Θ' is obtained by removing in Θ the formulas A and $A^\perp$ as well as our cut links and L, L' and adding two new cut links with respective premises $B, B^\perp$ and $C, C^\perp$ ($\otimes$ -reduction).
- If L is a $\&$ -link (with premises B and C) and L' is a $\oplus_1$ -link (with premise $B^\perp$), then Θ' is obtained in three steps : first we remove in Θ the formulas A and $A^\perp$ as well as our cut link and L, L' ; then we replace the eigenweight p_L by 1 and keep only those formulas and links that still have a nonzero weight : therefore B and $B^\perp$ remain with weight 1 whereas C disappears ; finally we add a cut between B and $B^\perp$ ($\oplus_1$ -reduction).
- If L is a $\&$ -link (with premises B and C) and L' is a $\oplus_2$ -link (with premise $C^\perp$), then Θ' is obtained in three steps : first we remove in Θ the formulas A and $A^\perp$ as well as our cut link and L, L' ; then we replace the eigenweight p_L by 0 and keep only those formulas and links that still have a nonzero weight : therefore C and $C^\perp$ remain with weight 1 whereas B disappears ; finally we add a cut between B and $B^\perp$ ($\oplus_2$ -reduction).

Proposition 1

If Θ' is obtained from a proof-net Θ by lazy cut-elimination, then Θ' is still a proof-net.

PROOF. — we consider all possible cut-elimination steps :

- *ID*-reduction : the only important thing to remark is that the other conclusion of L (which is an occurrence of $A^\perp$) cannot be the same occurrence (otherwise the structure would bear a cycle).
- $\otimes$ -reduction : let us fix a valuation φ , and let X be the set $\varphi(\Theta)$, that we can write as the union $Y \cup Z$ of $\varphi(\Theta')$ and $\{B, C, A, A^\perp, B^\perp, C^\perp\}$, so that $Y \cap Z = \{B, C, B^\perp, C^\perp\}$. Given a switching $\mathcal{S}$ of Θ' , it induces a graph $\Theta'_\mathcal{S}$ on Y and we can consider its subgraph $\mathcal{G}$ (on the same Y) in which the edges between $B, B^\perp$ and $C, C^\perp$ have been removed, as well as the subgraph $\mathcal{H}$ (on $Y \cap Z$) consisting of these two edges : we have $\Theta'_\mathcal{S} = \mathcal{G} \cup \mathcal{H}$. We want to show that $\Theta'_\mathcal{S}$ is connected and acyclic. In order to do this, we "extend" our switching $\mathcal{S}$ to a switching $\mathcal{S}'$ of Θ ; it is immediate that $\Theta_{\mathcal{S}'}$ can be written as $\mathcal{G} \cup \mathcal{H}'$, where $\mathcal{H}'$ is a graph on Z with one edge between B, A , one edge between C, A , one edge between $A, A^\perp$, one edge between $A^\perp, B^\perp$ (or between $A^\perp, C^\perp$, depending on the switching of L'). Since $\mathcal{G} \cup \mathcal{H}'$ is connected and acyclic, then $\mathcal{G}$ has 3 connected components, and each of them meets $Y \cap Z$. Now B, C are not in the same component (otherwise there would be a cycle in $\mathcal{G} \cup \mathcal{H}'$; $B, B^\perp$ are neither in the same component, since, if we switch L' to "left", we get a cycle in $\mathcal{G} \cup \mathcal{H}'$; the same is true, for symmetrical reasons, of $B, C^\perp$ and $C, C^\perp$. Then $B^\perp, C^\perp$ must lie in the same component. From this it is immediate that $\mathcal{G} \cup \mathcal{H}$, i.e. $\Theta'_\mathcal{S}$, is connected and acyclic.
- $(\oplus - 1)$ -reduction) : a valuation φ for Θ' can be extended to a valuation φ' for Θ by setting $\varphi'(p_L) = 1$ (and giving any value to the other eigenweights which might have disappeared when making the cut-elimination step). One can also extend a switching $\mathcal{S}$ of Θ' into a switching $\mathcal{S}'$ of Θ by setting $\mathcal{S}(L) = B$. Then it is almost immediate that $\Theta_{\mathcal{S}'}$ and $\Theta'_\mathcal{S}$ are of the same type.
- $(\oplus - 2)$ -reduction) : symmetrical to the previous case. □

Definition 15

The size of a proof-net is defined to be the number of its links.

Theorem 4

Lazy cut-elimination converges to a unique lazy normal form (i.e. a proof-net without ready cuts) in a time which is linear in the size of the proof-net.

PROOF. — unicity of the normal form is an easy consequence of the fact that our procedure is Church-Rosser. To prove termination, observe that some of the eigenweights receive definite values $p_L = 0$ or $p_L = 1$ during the cut-elimination of Θ . Let φ be a valuation of Θ which extends these choices. Now, if $\Theta \mapsto \Theta_1 \mapsto \dots \mapsto \Theta_n$ is a reduction sequence, we can observe the slices $\varphi(\Theta)$, $\varphi(\Theta_1)$, $\dots$, $\varphi(\Theta_n)$, and observe that the number of links in each of these slices is strictly decreasing. Then n cannot be bigger than the the number of links in $\varphi(\Theta)$, hence, n is smaller than the number of links in Θ . Lazy normalization therefore takes a linear number of steps, which is not quite the same as linear time. However, it is quite easy to see how to transform this into a linear time algorithm : we keep track of the values taken by the eigenweights and we delay the substitutions occurring in the additive reductions, i.e. we perform them only in case they yield values 0 (in which case we may erase) or 1 ... When the process is completed, then we perform the remaining substitutions. $\square$

Remark. — Only a limited part of the correctness criterion is actually used to prove proposition 1 ; the most important part is devoted to the absence of deadlock, i.e. theorem 3.

2 THE CASE OF QUANTIFIERS

Quantifiers have been treated in previous papers [G88, G90], hence we shall be brief. Roughly speaking the role of eigenweights is taken by eigenvariables, and everything is adapted *mutatis mutandis*.

2.1 Proof-structures

$$\begin{array}{llll} \forall \text{--links} & : & 1 \text{ premise} : A[e/x] & 1 \text{ conclusion} : \forall x A \\ \exists_t \text{--links} & : & 1 \text{ premise} : A[t/x] & 1 \text{ conclusion} : \exists x A \end{array}$$

$\forall$ -links make use of *eigenvariables* ; for each $\forall$ -link L there is a specific variable e_L which is associated with this link. Each existential link comes with the name of a term, namely the term t of the premise of the $\exists$ -link. This remark would be pure pedantism if we had not to take care of the case of fake dependencies (i.e. when x does not occur in A) where we cannot recover t from the premise. We shall say that a formula A *depends* on an eigenvariable e_L when

- either A is the premise of L
- or e_L occurs in A
- or A is the premise of a link $\exists_t$ and e_L occurs in t

Proof-structures are defined as expected i.e. as in definition 3 :

- if L is a quantifier link with premise A , then $w(L) = w(A)$;
- we require that $w(A) \cdot \neg w(L) = 0$ for any formula A depending on e_L ;
- we require that no eigenvariable occurs in a conclusion.

One easily defines what it means for a sequent calculus proof to be a sequentialisation of a given proof-structure, by extending definition 4 :

- If L is a $\forall$ -link with premise $A[e_L/x]$, and $\Gamma, \forall x A$ is the set of conclusions of Θ : the removal of L consists in removing the conclusion $\forall$ and the link L and replace everywhere e_L with a fresh variable e ; this induces a proof-structure with conclusions $\Gamma, A[e/x]$.
- If L is a $\exists_t$ -link with premise $A[t/x]$, and $\Gamma, \exists x A$ is the set of conclusions of Θ : the removal of L consists in removing the conclusion $\exists x A$ and the link L ; this induces a proof-structure with conclusions $\Gamma, A[t/x]$. Observe that this removal might be impossible, typically if some eigenvariable occurs in t .

2.2 Proof-nets

Definitions 8 and 9 are adapted as follows :

- $\mathcal{S}$ selects for each $\forall$ -link L of $\varphi_{\mathcal{S}}(\Theta)$ a formula $\mathcal{S}(L)$, the *jump* of L , depending on p_L in $\varphi_{\mathcal{S}}(\Theta)$. There is always a *normal* choice of jump for L , namely the premise A of L .
- In $\Theta_{\mathcal{S}}$ we draw an edge between the conclusion and the premise of any $\exists$ -link, and for all $\forall$ -links L of $\varphi_{\mathcal{S}}(\Theta)$, we draw an edge between the jump $\mathcal{S}(L)$ of L and the conclusion of L .

The condition for being a proof-net is exactly defined as in definition 10. We have only to check the extension of theorems 1 and 2. We only give some indications as to the proof of the latter :

- We shall mainly be concerned with a maximal empire eA where A is the conclusion of either a $\&$ or a $\forall$ -link.
- If eA is maximal, then its border formulas depend on no eigenvariable
- If eA is maximal and if $B, B' \in \varphi(\Theta)$ are such that e'_L occurs in B, B' ; then $B \in eA$ implies $B' \in eA$.

2.3 Cut-elimination in MALLq

Lazy cut-elimination is defined by adding a case to definition 14 :

- If L is a $\forall$ -link (with premises $B[e_L/x]$) and L' is a $\exists_t$ -link (with premise $B^\perp[t/x]$), then Θ' is obtained in three steps : first we remove in Θ the formulas A and $A^\perp$ as well as our cut link and L, L' ; then we replace the eigenvariable p_L by t ; finally we add a cut between $B[t/x]$ and $B^\perp[t/x]$ ($\exists_t$ -reduction).

The procedure might not work for general proof-structures : if e_L actually occurs in t , then substituting t for e_L changes $B^\perp[t/x]$ into $B^\perp[t'/x]$ which does not match $B[t/x]$, and also e_L is still present. But the proof-net condition forbids this (just set $\mathcal{S}$ so that $\mathcal{S}(L) = B^\perp[t/x]$).

It is of course essential to prove that the proof-net condition is preserved through cut-elimination, which essentially amounts looking at our new step : given a switching $\mathcal{S}$ of Θ' , we want to back it up into a switching $\mathcal{S}'$ of Θ , which offers no difficulty as to L (set $\mathcal{S}'(L) = B[e_L/x]$) ; but this might be problematic for some $\forall$ -link L' , since a formula C may not depend on e'_L , whereas $C[t/e_L]$ does, and if $\mathcal{S}(L') = C[t/e_L]$, we cannot set $\mathcal{S}'(L') = C$. However observe that :

- in such a case, e'_L must occur in t , and this forces (w.r.t. Θ) $B[t/x] \in eA'$, hence $\forall x A \in eA'$, where A' is the conclusion of L' . But since $B[t/x] \notin eA$, we also get $A' \notin eA$, hence $eA \subset eA'$
- C must also depend on e_L , hence $C \in eA$, which implies $C \in eA'$
- but it is immediate that the illegal jump $\mathcal{S}'(L') = C$ to an element of eA' will not alter the correctness criterion.

We can therefore prove that the proof-net condition is preserved by induction on the number of illegal jumps (with an induction loading : we consider proof-nets in which fake dependencies have been declared). The basis step being trivial, the induction step ($n + 1$ illegal jumps) consists in choosing L', A', C as above, and to modify our proof-net by introducing a fake dependency i.e. to pretend that e'_L occurs in C . This induces another proof-net (as observed above) in which our jump is now legal. There are only n illegal jumps in this new proof-net, and the induction hypothesis applies, yielding a connected acyclic graph.

2.4 Complexity of cut-elimination

As before, cut-elimination involves a linear number of steps. This is clearly not enough for linear time normalization, since one of these steps is a substitution, whose iteration is exponential. However, the algorithm remains linear

if we do not perform the substitutions, i.e. we keep the eigenvariables in the proof-net and we use an auxiliary stack to remember which term should be substituted for them. This process is particularly efficient in the following case : we normalize a proof of A without $\&$ or $\exists$; then the lazy cut-elimination will eventually yield a cut-free proof (in which some substitutions have not been performed) in linear time. Now, imagine that we perform the delayed substitutions in a given formula B of our net : we get a subformula B' of A which is completely determined by the choice of the links below B , hence we can forget the substitutions and directly write this subformula.

3 EXPONENTIALS

We shall deal with only two rules, which are central in the study of exponentials, namely weakening and contraction. The point of the introduction of exponentials is precisely to allow weakening and contraction for certain formulas, those which are prefixed by $?$. This is not enough, and in the standard version of linear logic, two additional rules, promotion and dereliction, are added. We now know that the choice of additional rules is much more open, and we shall therefore ignore them. As a consequence, it will be impossible to speak of cut-elimination, and we shall concentrate on correctness.

Our basic ingredient will be the notion of a *discharged* formula (terminology strictly inspired from natural deduction). Besides usual formulas we shall allow in a proof-net *discharged* ones, denoted $[A]$. These formulas are handled in a very specific way :

- If $[A]$ is the premise of L , then L is a $?$ -link (see below).
- If $[A]$ is a conclusion of the proof-net, we no longer require its weight to be 1. We therefore modify definition 10 as follows : a non-discharged conclusion has weight 1.
- We shall assume that $[A]$ is the conclusion of (unspecified) generalized axioms (i.e. boxes). This will handle all the additional rules we know, but usual dereliction³.
- Since discharged formulas are bound to be merged by $?$ -links (see below), there is no need to superimpose such formulas, hence we require that discharged formulas are the conclusions of a unique link.

The only link we consider is the n -ary link $?_n$, with n unordered premises, which are all occurrences of the same discharged formula $[A]$, and one conclusion,

3. Usual dereliction easily fits into our pattern : just allow any formula which is the conclusion of a link to be discharged.

namely $?A$. The case $n = 0$ is allowed, and accounts for weakening. The condition for being a proof-structure is the following : the weight increases, i.e. if $[A_i]$ is any premise of L , then $w([A_i]) \leq w(L)$. The condition for being a proof-net depends on an additional datum : for any $?-link$ L , a default jump jL is chosen : jL can be any formula B (discharged or not) such that $w(L) \leq w(B)$. Given a valuation φ , a switching $\mathcal{S}$ will select a jump for each $?-link$ L : this jump may be the default one jL or any premise $[A_i]$ of L such that $\varphi(w([A_i])) = 1$. We then state the usual connected-acyclic condition. One must introduce a sequent calculus with two kinds of formulas, usual ones and discharged ones. This calculus is identical to the usual but for :

- the $\&$ -rule : the contexts may differ on some discharged formulas. From $\vdash \Gamma, [\Delta], A$ and $\vdash \Gamma, [\Delta'], B$, deduce $\vdash \Gamma, [\Delta], [\Delta'], A \& B$.
- the new rule of weakening/contraction : from $\vdash \Gamma, [A], \dots, [A]$, deduce $\vdash \Gamma, ?A$.

Sequentialisation is an easy exercise.

A APPENDIX

A.1 More about additives

A.1.1 An alternative to weights

Let Θ be a proof-net ; then we can define a structure of coherent space, on the set consisting of all formulas and all links of Θ : $X \subset Y$ iff $w(X).w(Y) \neq 0$. Now, due to the fact that weights are monomials, as soon as $X_1, \dots, X_n$ are pairwise coherent, the intersection $w(X_1) \cap \dots \cap w(X_n)$ is nonzero. This means that maximal cliques in the coherent space correspond to slices. This also means that we can replace the weighting of Θ by the coherent space structure⁴. This alternative presentation has many advantages, in particular that we have not to name the eigenweights. This becomes crucial with the technique of spreading that we now introduce, since a spreading may introduce new eigenweights.

A.1.2 Spreading of a proof-net

Let Θ be a proof-net, let A be a formula in Θ and let p be an eigenweight. The spreading of Θ above A w.r.t. p consists in replacing every formula or link hereditarily above A (including A) and whose weight does not depend on p by two copies, X_1 and X_2 . This induces a new coherent space structure :

- $X_i \subset Y_j$ iff $i = j$ and $X \subset Y$ in Θ

4. In this way the dependency condition looks very natural

- $X_1 \subset Y$ iff $w(X).w(Y).p \neq 0$ in Θ
- $X_2 \subset Y$ iff $w(X).w(Y).\neg p \neq 0$ in Θ
- $X \subset Y$ iff $X \subset Y$ in Θ

The resulting coherent space needs not be a proof-structure. However, in the case we shall use this construction, this will be the case, and the resulting coherent space will indeed be a proof-net.

A.1.3 Cut-elimination

Let us say that a cut is *almost ready*, when both premises are the conclusions of a unique link. For almost ready cuts there is an obvious cut-elimination, the same as in the ready case. We shall now explain how to reduce a general cut to almost ready ones.

Assume that the cut is between A and $A^\perp$, with the same monomial weight w ; let $w_1, \dots, w_m$ be the weights of the links with conclusion A , and $w'_1, \dots, w'_n$ the weights of the links with conclusion $A^\perp$, so that

$w = w_1 + \dots + w_m = w'_1 + \dots + w'_n$; we assume that the cut is not at most ready, so $m + n > 2$. Let us say that an eigenweight p *splits* A when we can partition $w_1, \dots, w_m$ into two non-empty subsets such that the partial sums are respectively equal to $w.p$ and $w.\neg p$; as soon as $m > 1$ there is at least one splitting for A .

- if the eigenweight p is a common splitting for A and $A^\perp$, then we can duplicate A into two copies of respective weights $w.p$ and $w.\neg p$, the same for $A^\perp$, and replace our cut with two cuts
- otherwise, there is a splitting, of say, $A^\perp$ w.r.t. an eigenweight p ; we spread Θ above A w.r.t. p and we are back to the previous case. (The spreading makes sense mainly because $A^\perp \in eA$ w.r.t. any valuation, hence that all formulas above $A^\perp$ are in eA).

This procedure is Church-Rosser and terminating. It should be considered when theoretical questions (subformula property etc.) are at stake. Its computational value is limited, since iterated spreadings may induce exponentially many duplications.

A.1.4 Discussion

We would like to discuss several issues connected with additive proof-nets. Contrarily to the multiplicative case, the extant solution is not perfect (although it has the virtue to exist). Let us discuss the weaknesses of our solution and potential improvements :

the main question which occurs when translating sequent calculus to proof-nets is the problem of superimposition. There are two difficulties :

A.1.5 Weights

Weights must be monomials. However, weights of the form $p \cup q$ will naturally occur if we want to allow more superimpositions. The present state of affairs is as follows :

- ▶ in spite of years of efforts, I never succeeded in finding the right correctness criterion for these more liberal proof-nets
- ▶ general boolean coefficients might be delicate to represent (on the other hand, the case we consider has a natural presentation in terms of coherent spaces)
- ▶ normalization in the full case might be messy

A.1.6 Identity of formulas

Imagine that we have no problem with weights, and that we try to maximize the identifications. One idea is to adopt a binary $\oplus$ -rule, which is quite natural. Besides that, we immediately stumble on a lot of problems :

- ▶ if I have a cut on A in both proofs, should I superimpose them... worse, if both proofs have two cuts on A , which ones should be superimposed ? The same question occurs with the contraction rule : how do we superimpose two contraction rules (remember that the premises are indiscernible).
- ▶ what about existential witnesses ? Should one superimpose two portions of proofs with the same structure, but different witnesses ?
- ▶ what about normalization ? During this process, distinct portions of proofs might become equal, hence it would be necessary to superimpose them...

These limitations do not apply if we restrict to cut-free proofs in **MALL**, and if we make the extra assumption that all identity links are atomic : there is a well defined notion of net (provided one can fix the problem of weights) which enjoys the maximum number of identifications.

If identification is difficult, its converse is easy, i.e. there is not the slightest problem to forget that two formulas are equal. This could indicate a possible theoretical way out, namely considering a proof-net as the set of its slices. The computational value of this idea is limited (exponential growth of the net), but this might be valuable for theoretical considerations. However, we have no idea how to define a correctness condition for such sets of slices.

There are other problems connected to normalization :

A.1.7 Normalization

Multiplicative proof-nets have a local cut-elimination, and this is still true for quantifiers, provided we do not compute the existential witnesses. The additive case involves a real global move, which consists in setting an eigenweight to 0 or 1, and erase everything with weight 0. This is rather brutal, and completely foreign to the parallel asynchronous spirit of proof-nets. In [G94] (section 3.4.) a variant of usual sequent calculus is introduced, with a local cut-elimination, i.e. the erasing is performed in a lazy way, which means that some useless "beards", which are bound to be erased, are still hanging. The calculus with these beards is connected with the additive neutral $\top$ (which has still no satisfactory treatment), and we can expect that there is a notion of additive bearded proof-net with a local elimination. This is perhaps the most promising open question in this paper.

A.2 Multiplicative neutrals

There are two multiplicative neutrals, 1 and $\perp$, and two rules, the axiom $\vdash 1$ and the weakening rule : from $\vdash \Gamma$, deduce $\vdash \Gamma, \perp$. Both rules are handled by means of links with one conclusion and no premise ; however $\perp$ -links are treated like 0-ary ?-links, i.e. they must be given a default jump. Sequentialisation is immediate.

At first sight, cut-elimination is unproblematic : replace a cut between the conclusions 1 and $\perp$ of zero-ary links with... nothing. But we notice a new problem, namely that a cut formula A can be the default jump of a $\perp$ -link L , and we must therefore propose another jump for L . Usually one of the premises of the link with conclusion A works (or the jump of L' if A is the conclusion of a $\perp$ -link) works. Worse, this new jump is by no ways natural (if A is $B \otimes C$, the new jump can either be B or C), which is quite unpleasant. As far as we know, the only solution consists in declaring that the jumps are not part of the proof-net, but rather of some control structure. It is then enough to show that at least one choice of default jump is possible. This is not a very elegant solution : we are indeed working with equivalence classes of proof-nets and if we want to be rigorous we shall have to endlessly check that such and such operation does not depend on the choice of default jumps. In practice one can be rather sloppy...

Of course everything would be nicer without any default jump. But then a proof-net for a multiplicative combination A of occurrences of 1 and $\perp$ would basically be nothing more than A itself : the correctness criterion for proof-nets without jumps encompasses the decision problem for such combinations, and this problem is known to be NP-complete by [LW92]... The existence of a correctness criterion of the same style as the familiar ones is therefore very

unlikely⁵.

This discussion is fully relevant to the exponential case : as soon as we start to normalize exponential cuts, then the same problems (with the same solution) arise.

A.3 Additive neutrals

There is still no satisfactory approach to additive neutrals, which are fortunately extremely uninteresting in practice. The only way of handling $\top$ is by means of a box or, if one prefers, by means of a second order translation : on this Kamtchatka of linear logic, the old problems of sequent calculus are not fixed. The absence of a satisfactory treatment of $\top$ calls for another notion of proof-net... presumably a solution to the wider question of bearded proof-nets.

BIBLIOGRAPHY

- [DR88] **Danos, V. & Regnier, L. :** The structure of multiplicatives, *Archive for Mathematical Logic* 28.3, pp. 181-203, 1989
- [G86] **Girard, J.-Y. :** Linear Logic, *Theoretical Computer Science* 50.1, pp. 1-102, 1987
- [G86A] **Girard, J.-Y. :** Multiplicatives, *Rendiconti del Seminario Matematico dell'Università e Policlinico di Torino, special issue on Logic and Computer Science*, pp. 11-33, 1987
- [G87A] **Girard, J.-Y. :** Towards a Geometry of Interaction, *Categories in Computer Science and Logic, Contemporary Mathematics* 92, AMS 1989, pp. 69-108
- [G88] **Girard, J.-Y. :** Quantifiers in linear logic, *Temi e prospettive della logica e della filosofia della scienza contemporanee*, CLUEB, Bologna, 1989.
- [G90] **Girard, J.-Y. :** Quantifiers in linear logic II, *Nuovi problemi della logica e della filosofia della scienza*, CLUEB Bologna 1991.

5. There is still a possibility, namely to work with the necessary condition *acyclic graph with $n + 1$ connected components*, where n is the number of $\perp$ -links ; but we would have to check that this condition is preserved through cut-elimination (which seems likely), and we would have to restrict to proof-nets in which the formula $\perp$ cannot occur in the absence of cuts, so that our necessary condition eventually becomes sufficient...

- [G94] **Girard, J.-Y.** : Geometry of interaction III : the general case, *Proceedings of the Workshop on Linear Logic*, MIT Press, submitted.

- [LMSS90] **Lincoln, P. & Mitchell, J. & Scedrov, A. & Shankar N.** : Decision Problems in Propositional Linear Logic, *Annals of Pure and Applied Logic* 56,1992, pp. 239-311.

- [LW92] **Lincoln, P. & Winkler, T.** : Constant-only Multiplicative Linear Logic is NP-complete, *Theoretical Computer Science* 135,1994,pp.155-159.