

Sequentialization of Multiplicative Proof Nets

Olivier . Laurent @ens-lyon.fr

April 2, 2013

We present a new simple proof of the sequentialization property for proof nets of unit-free multiplicative linear logic [Gir87] satisfying the Danos–Regnier correctness criterion [DR89].

Multiplicative Proof Nets

We recall some basic definitions, terminology, ... about MLL proof nets and the Danos–Regnier correctness criterion.

Proof structures. A *proof structure* $\mathcal{S}$ is a non-empty finite directed acyclic multigraph with loose edges, that is the data $(\mathcal{N}, \mathcal{C}, \mathcal{E}, \mathbf{s}, \mathbf{t})$ where $\mathcal{N}$ is the set of *nodes*, $\mathcal{C}$ is the set of *sinks* (these are the two kinds of vertices), $\mathcal{E}$ is the set of *edges*, $\mathbf{s}$ is the *source* function from $\mathcal{E}$ to $\mathcal{N}$ and $\mathbf{t}$ is the *target* function from $\mathcal{E}$ to $\mathcal{N} + \mathcal{C}$, such that each element of $\mathcal{C}$ is the image through $\mathbf{t}$ of exactly one edge, and such that $(\mathcal{N} + \mathcal{C}, \mathcal{E}, \mathbf{s}, \mathbf{t})$ is a non-empty finite directed acyclic multigraph.

If $\mathbf{s}(e) = N$ we call e a *conclusion* of N , if $\mathbf{t}(e) = N$ we call e a *premise* of N , if $\mathbf{t}(e) \in \mathcal{C}$ (i.e. e is a loose edge) we call e a *conclusion* of $\mathcal{S}$. If all the conclusions of a node N are conclusions of $\mathcal{S}$, we call N a *terminal* node.

In a proof structure, each node is labelled with its kind ax , $\otimes$ or $\wp$ and edges are labelled with formulas of MLL (this label is also called the *type* of the edge). It is required that:

- Each node labelled ax has no premise and two conclusions with dual types A and $A^\perp$.
- Each node labelled $\otimes$ has two premises which are ordered and one conclusion. This way we can speak about the *left* (first) and *right* (second) premise of a $\otimes$ node. If the type of the left premise is A and the type of the right premise is B , the type of the conclusion must be $A \otimes B$.
- Each node labelled $\wp$ has two premises which are ordered and one conclusion. If the type of the left premise is A and the type of the right premise is B , the type of the conclusion must be $A \wp B$.

We consider the cut-free case only, since as usual cuts can be encoded as $\otimes$ nodes for sequentialization in MLL.

Switchings. Given a proof structure $\mathcal{S}$, let $\mathcal{P}$ be the set of its $\wp$ nodes. For each function $\varphi : \mathcal{P} \rightarrow \{\text{left}, \text{right}\}$ (called a *switching*), we define the *correctness graph* $\mathcal{S}_\varphi$ by turning the $\varphi(P)$ premise of each $\wp$ node P into a new loose edge (where φ exchanges *left* and *right* in the values of φ). Formally, $\mathcal{S}_\varphi$ is the non-empty finite directed acyclic multigraph with loose edges

$(\mathcal{N}, \mathcal{C} + \mathcal{C}_{\mathcal{P}}, \mathcal{E}, \mathbf{s}, \mathbf{t}_{\varphi})$ where $\mathcal{C}_{\mathcal{P}} = \{C_P \mid P \in \mathcal{P}\}$ (C_P is a new sink associated with each P) and:

$$\mathbf{t}_{\varphi}(e) = \begin{cases} \mathbf{t}(e) & \text{if } \mathbf{t}(e) \notin \mathcal{P} \\ P = \mathbf{t}(e) & \text{if } e \text{ is the } \varphi(P) \text{ premise of } P \\ C_P & \text{if } e \text{ is the } \overline{\varphi}(P) \text{ premise of } P \end{cases}$$

$\mathcal{S}_{\varphi}$ is not a proof structure: $\mathfrak{V}$ nodes have only one premise.

Danos–Regnier correctness. A proof structure $\mathcal{S}$ is *DR-correct* if for each correctness graph $\mathcal{S}_{\varphi}$, its undirected underlying graph is acyclic and connected.

Fact 1. Any DR-correct proof structure is a (weakly) connected directed multigraph.

Sequentialization

Sequentialization Process. Let $\mathcal{S}$ be a proof structure and N be a terminal node (of kind $\otimes$ or $\mathfrak{V}$). The *removal* of N in $\mathcal{S}$ is the proof structure obtained by removing N , the conclusion of N (if any) with its target and, by adding two new sinks as targets for the premises of N which are now loose edges. If $\mathcal{S}$ is DR-correct, a terminal node N is said to be *sequentializing* if, depending on its kind:

- *ax* node: N is the only node of $\mathcal{S}$.
- $\otimes$ node: the removal of N in $\mathcal{S}$ has two (weakly) connected components being DR-correct proof structures.
- $\mathfrak{V}$ node: the removal of N in $\mathcal{S}$ is a DR-correct proof structure.

Given a DR-correct proof structure $\mathcal{S}$, the *sequentialization process* builds a sequent calculus proof as follows. Assume $\mathcal{S}$ contains a sequentializing terminal node N , we look at the kind of N :

- *ax* node: we stop successfully with the proof reduced to an (*ax*) rule.
- $\otimes$ node: we apply recursively the sequentialization process to the two obtained DR-correct proof structures (in the removal of N). We apply a ($\otimes$) rule to the two obtained proofs.
- $\mathfrak{V}$ node: we apply the sequentialization process to the removal of N , and we apply a ($\mathfrak{V}$) rule to the obtained proof.

A proof structure $\mathcal{S}$ is called *sequentializable* if the sequentialization process succeeds, that is if it is possible to find a sequentializing terminal node at each step.

To prove that DR-correct proof structures are sequentializable, it is thus enough to show that any DR-correct proof structure has a sequentializing terminal node. This is what we do now.

Paths. A (undirected) *path* γ is a finite sequence of pairs $(e_i, \epsilon_i)_{1 \leq i \leq n}$ ($n \in \mathbb{N}$) with $e_i \in \mathcal{E}$ and $\epsilon_i \in \{+, -\}$ (its *sign*) such that $\mathbf{t}_{\gamma}(e_i) = \mathbf{s}_{\gamma}(e_{i+1})$ ($1 \leq i < n$), where the γ -*source* $\mathbf{s}_{\gamma}(e_i)$ is $\mathbf{s}(e_i)$ if $\epsilon_i = +$ and $\mathbf{t}(e_i)$ if $\epsilon_i = -$, and the γ -*target* $\mathbf{t}_{\gamma}(e_i)$ is $\mathbf{t}(e_i)$ if $\epsilon_i = +$ and $\mathbf{s}(e_i)$ if $\epsilon_i = -$. The *source* (resp. *target*) of γ is $\mathbf{s}(\gamma) = \mathbf{s}_{\gamma}(e_1)$ (resp. $\mathbf{t}(\gamma) = \mathbf{t}_{\gamma}(e_n)$). A path is *simple* if it never goes twice through the same node: $\mathbf{t}_{\gamma}(e_i) = \mathbf{s}_{\gamma}(e_j) \Rightarrow j = i + 1$ for $1 \leq i < j \leq n$ (we allow the case of a cycle: $\mathbf{t}_{\gamma}(e_n) = \mathbf{s}_{\gamma}(e_j)$ or $\mathbf{t}_{\gamma}(e_j) = \mathbf{s}_{\gamma}(e_1)$ for some $1 \leq j \leq n$). If N_1 and N_2 are

two occurrences of nodes visited by the path γ , we denote by $\gamma_{N_1 N_2}$ the sub-path of γ going from N_1 to N_2 . The path $\bar{\gamma}$ is the *reverse* of γ (edges in reverse order with opposite signs): $\bar{\gamma} = (e_{n+1-i}, -\epsilon_{n+1-i})_{1 \leq i \leq n}$. If γ_1 and γ_2 are two paths such that $t(\gamma_1) = s(\gamma_2)$ (we call them *compatible*) then their concatenation $\gamma_1 \cdot \gamma_2$ is a path.

A *locally switching* path is a simple path which does not start from a $\mathfrak{V}$ node through one of its premises ($s_\gamma(e_1)$ is a $\mathfrak{V}$ node implies $\epsilon_1 = +$) and does not go consecutively through the two premises of a $\mathfrak{V}$ node (if $t_\gamma(e_j)$ is a $\mathfrak{V}$ node for some $1 \leq j < n$ then $\epsilon_j = -$ or $\epsilon_{j+1} = +$).

Fact 2. *If γ_1 and γ_2 are two compatible disjoint non-cyclic locally switching paths, then $\gamma_1 \cdot \gamma_2$ is a locally switching path.*

Since edges of a proof structure $\mathcal{S}$ and of its correctness graphs $\mathcal{S}_\varphi$ are the same, it is meaningful to compare their paths. A path in $\mathcal{S}_\varphi$ is always a path in $\mathcal{S}$ (but the converse is wrong in general). Such a simple path of $\mathcal{S}$ which is a path in some $\mathcal{S}_\varphi$ is called a *globally switching* path. DR-correctness implies the absence of globally switching cycle.

Fact 3. *A globally switching path, which does not start with the premise of a $\mathfrak{V}$ node, is a locally switching path.*

Lemma 1 (Local-global principle)

A locally switching cycle γ induces a globally switching cycle as soon as it does not end with the premise of a $\mathfrak{V}$ node P such that the other premise e of P satisfies $(e, -) \in \gamma$.

PROOF: If a cycle does not contain the two premises of any $\mathfrak{V}$ node, it is a globally switching cycle. Let γ' be the minimal sub-path of γ which is a cycle. Since γ is a locally switching cycle, if γ' does not start or end with the premise of $\mathfrak{V}$ node, it is then a globally switching cycle. If γ' starts and ends with the premises of the same $\mathfrak{V}$ node, since γ does not start with such a premise, γ' must be a suffix of γ . We then have a contradiction with the hypotheses. $\square$

If e is an edge, its *descending path* $\delta(e)$ is the unique directed path (*i.e.* using only $+$ signs) starting from e and ending with the premise of a terminal node, if it exists (otherwise $\delta(e)$ is empty). If N is a $\mathfrak{V}$ or $\otimes$ node, we note $\delta(N)$ the descending path of its unique conclusion. $\delta(N)$ is empty if and only if N is terminal.

Fact 4. *$\delta(N)$ is always a locally switching path.*

Correctness $\mathfrak{V}$ node. Let T be a $\otimes$ node, a *correctness $\mathfrak{V}$* for T is a $\mathfrak{V}$ node P with two disjoint locally switching paths κ_0 and κ_1 from T to P (called *correctness paths*) which both start with a premise of T and end with a premise of P . A triple (κ_0, κ_1, P) is called a *correctness triple* for T and a pair (κ_i, P) ($i \in \{0, 1\}$) a *correctness pair*.

Lemma 2 (Correctness $\mathfrak{V}$)

Non-sequentializing terminal $\otimes$ nodes of DR-correct proof structures have correctness $\mathfrak{V}$ s.

PROOF: Let T be a terminal $\otimes$ node and φ be a switching, the removal of T splits $\mathcal{S}_\varphi$ into two connected components (thanks to DR-correctness). If all $\mathfrak{V}$ nodes are such that both their premises belong to the same connected component, then T is a sequentializing $\otimes$ node of $\mathcal{S}$ since the removal of T in $\mathcal{S}$ has two connected components as well (which are DR-correct). Otherwise there exists a $\mathfrak{V}$ node P with a premise in each connected component of the removal of T in $\mathcal{S}_\varphi$. Each of these components contains a premise of T and a premise of P and (by connectivity) a path from the first to the second. The two obtained paths

are locally switching (Fact 3) and disjoint. Finally, in the component containing P , the obtained path cannot contain the conclusion of P otherwise, one could connect the two paths and obtain a cycle in the correctness graph $\mathcal{S}_{\varphi'}$ where φ' is obtained from φ by just changing the value for P . $\square$

Core Part of the Proof. We now fix a DR-correct proof structure $\mathcal{S}$. We are going to build a path leading to a sequentializing node.

$\mathcal{S}$ contains at least one ax node A . If A is a terminal node then it is the unique node of $\mathcal{S}$ (Fact 1) and it is sequentializing, we are done. Otherwise let e be a conclusion of A which is not a conclusion of $\mathcal{S}$, we follow $\delta(e)$ and we reach a terminal node N . $(\star)$ If N is a $\mathfrak{V}$ node, it is sequentializing. Otherwise N is a $\otimes$ node. Either it is a sequentializing $\otimes$ node, or we can apply Lemma 2 to obtain a correctness triple (κ, κ', N') . Starting from N' , we then follow $\delta(N')$ and we reach a terminal node. We can now continue as before from $(\star)$.

If we stop, it means we reach a sequentializing node and we are done. Otherwise we build an infinite sequence of (signed) edges $\mu = \kappa_1 \cdot \delta_1 \cdot \kappa_2 \cdot \delta_2 \cdot \kappa_3 \cdot \delta_3 \cdots$ where if $N_i = s(\kappa_i)$ and $N'_i = t(\kappa_i)$ then (κ_i, N'_i) is a correctness pair for N_i (which can be extended into a correctness triple $(\kappa_i, \kappa'_i, N'_i)$) and $\delta_i = \delta(N'_i)$ (thus N_i is a $\otimes$ node and N'_i is a $\mathfrak{V}$ node). Since $\mathcal{S}$ is finite, μ must visit twice the same node, and we are going to show this is not possible.

Let N be the first node in μ (thus N belongs to some κ_j or δ_j) which also has another occurrence in some κ_i or κ'_i or δ_i for some $i \leq j$ (note the crucial introduction of κ'_i here). If we are in one of the first two cases (N occurs in κ_i or κ'_i), we note $\kappa = \kappa^1 \cdot \kappa^2$ the element of $\{\kappa_i, \kappa'_i\}$ in which N occurs (and κ' the other element), with κ^1 ending with N and κ^2 starting with N . If κ^1 ends with the conclusion of N and κ^2 starts with one of its premises, we says N occurs badly (otherwise it occurs well). We define the path μ' as:

$$\mu' = \begin{cases} \delta' \cdot \mu_{N_{i+1}N} & \text{if } N \text{ occurs in } \delta_i \text{ at the beginning of its suffix } \delta' \\ \kappa^2 \cdot \mu_{N'_iN} & \text{if } N \text{ occurs well in } \kappa_i \text{ or } \kappa'_i \\ \kappa' \cdot \mu_{N'_iN} \cdot \overline{\kappa^1} & \text{if } N \text{ occurs badly in } \kappa_i \text{ or } \kappa'_i \end{cases}$$

where $\mu_{N'_iN}$ (resp. $\mu_{N_{i+1}N}$) is the sub-path of μ going from N'_i (resp. N_{i+1}) to the second visit of N by μ .

μ' is a locally switching path (Facts 4 and 2) which is a cycle and satisfies the hypotheses of Lemma 1. We thus have a globally switching cycle, this contradicts DR-correctness.

Figure 1 represents the key ingredients of the proof above (in the case where $\kappa' = \kappa_i$ and N occurs badly).

References

- [DR89] Vincent Danos and Laurent Regnier. The structure of multiplicatives. *Archive for Mathematical Logic*, 28:181–203, 1989.
- [Gir87] Jean-Yves Girard. Linear logic. *Theoretical Computer Science*, 50:1–102, 1987.

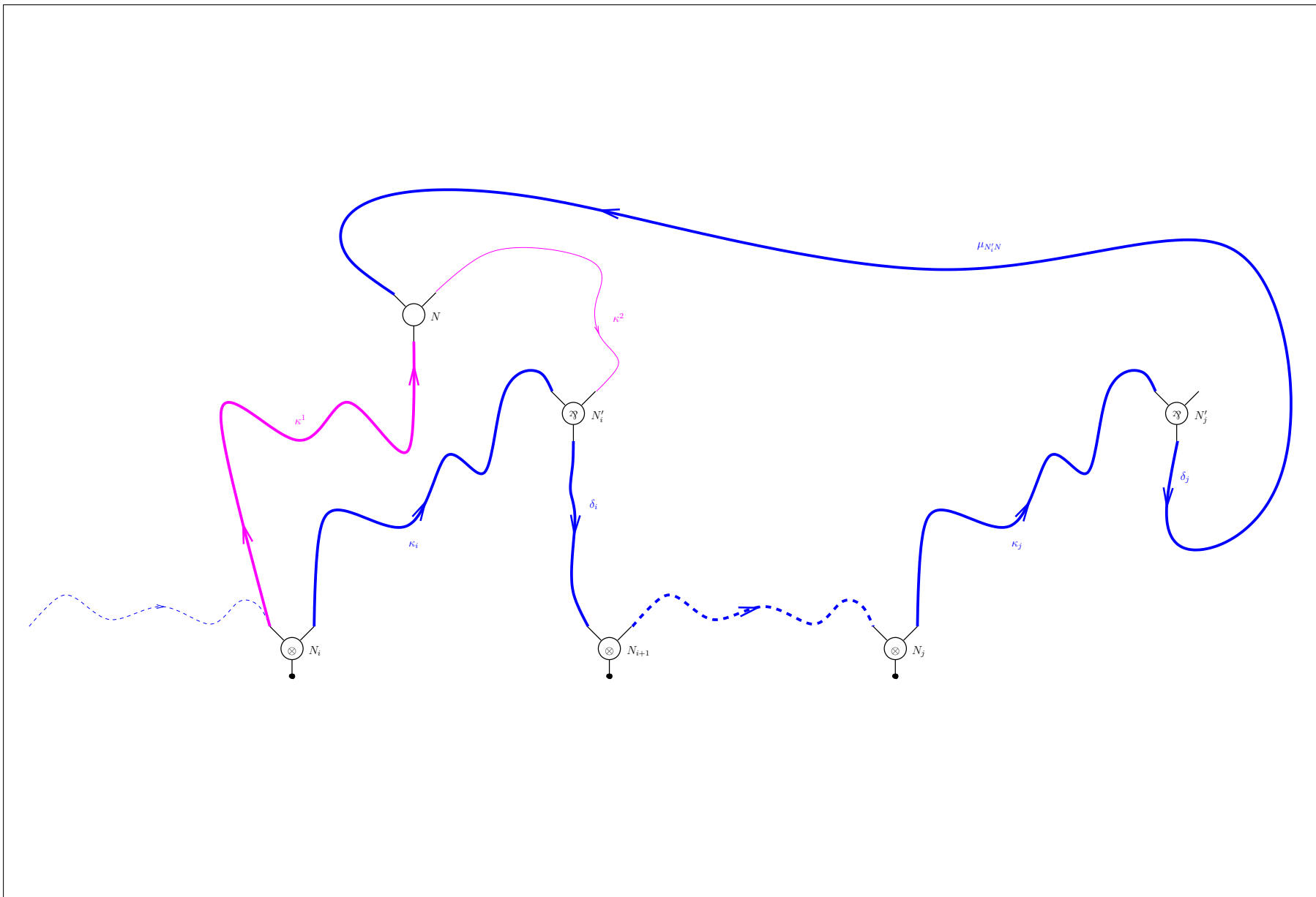


Figure 1: Cycle for $\kappa' = \kappa_i$