

Modular Constructive Lyndon Interpolation for Nondistributive Logics

Andrea De Domenico^{1,*†}, Giuseppe Greco^{2†}, Alessandra Palmigiano^{2,3†} and Apostolos Tzimoulis^{4†}

¹IMDEA Software Institute, Madrid, Spain

²Vrije Universiteit Amsterdam, The Netherlands

³Department of Mathematics and Applied Mathematics, University of Johannesburg, South Africa

⁴Department of Computer Science, University of Luxembourg

Abstract

We establish the Lyndon interpolation property for basic lattice expansion logics (LE-logics) in arbitrary signatures using display calculi. Our approach is *constructive*, yielding interpolants algorithmically from derivations, and *modular*, in the sense that interpolation for axiomatic extensions can be obtained by verifying a local interpolation property for the analytic structural rules corresponding to the additional axioms. To this end, we identify a class of *interpolation-safe* structural rules preserving local Lyndon interpolation. As applications of the general framework, we show that the tense version of Holliday’s fundamental modal logic enjoys the Lyndon interpolation property.

Keywords

Craig interpolation, Lyndon interpolation, Non classical logics, Lattice expansion logics, Display calculi

1. Introduction

In its classical form, the *Craig interpolation property* states that whenever a formula $A \rightarrow B$ is derivable in a logic, there exists a formula C , called an *interpolant*, such that both $A \rightarrow C$ and $C \rightarrow B$ are derivable, and moreover C contains only the non-logical symbols common to A and B . More generally, in abstract formulations of interpolation, implication is replaced by the derivability relation of the logic, while the formulas A and B may be replaced by contexts or bunches of formulas, either in the language of the logic itself or in a suitable extension thereof. The interpolant C , however, remains a formula of the original language of the logic, possibly occurring within a non-empty context. The stronger *Lyndon interpolation property* additionally requires preservation of the *polarity* of propositional variables. Interpolation has deep connections with algebra, model theory, proof theory, and category theory, as well as important applications in computer science, including modular specification, formal verification, model checking, and automated reasoning. In algebraic logic, interpolation is closely related to amalgamation and superamalgamation properties, while explicitly constructing interpolants usually requires the use of analytic calculi with good structural properties.

The study of interpolation in non-classical logics has produced a broad and technically diverse literature. Surveys such as [1] and [2] document the variety of semantic, algebraic, and proof-theoretic techniques developed for various modal and substructural logics. Among proof-theoretic approaches, a particularly fruitful line of research investigates interpolation via analytic calculi from which interpolants can be extracted algorithmically [3, 4].

CI-BD-SOQE 2026: Workshop on Craig Interpolation, Beth Definability, and Second-Order Quantifier Elimination, at FLoC 2026: The 9th Federated Logic Conference, July 24–25, 2026, Lisbon, Portugal

*Corresponding author.

†These authors contributed equally.

✉ andrea.domenico@imdea.org (A. De Domenico); g.greco@vu.nl (G. Greco); alessandra.palmigiano@vu.nl (A. Palmigiano); apostolos@tzimoulis.eu (A. Tzimoulis)

ORCID 0000-0002-8973-7011 (A. De Domenico); 0000-0002-4845-3821 (G. Greco); 0000-0001-9656-7527 (A. Palmigiano); 0000-0002-6228-4198 (A. Tzimoulis)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

The results of the present paper are closest to Kuznets’ modular approach to interpolation based on *labelled sequent calculi* [5], which adapts results obtained for a class of *internal* sequent-like formalisms [6]. In this line of research, interpolants are constructed inductively from derivations in analytic labelled calculi, yielding constructive proofs of Craig and Lyndon interpolation. One of the key features of this method is its modularity: interpolation for axiomatic extensions can be established by verifying local conditions on the additional rules corresponding to frame conditions. In particular, Kuznets identified a broad class of rules for which interpolation follows uniformly: *Horn-like geometric rules*, which equivalently correspond to frame conditions that can be captured by quantifier-free Horn formulas. This programme has significantly expanded the scope of proof-theoretic interpolation methods for substructural modal logics.

The present paper develops this modular and constructive perspective in the setting of lattice expansion logics (*LE-logics* [7, 8]) and their proper display calculi [9]. A characteristic feature of LE-logics is that the distributive laws of conjunction and disjunction are not necessarily valid, while the settings of most existing proof-theoretic results on interpolation assume the validity of distributive laws. This nondistributive setting has recently attracted increasing attention [10, 11, 12].

Proper display calculi provide a particularly natural framework for this investigation, because they enjoy fundamental properties such as the subformula property and cut elimination in a modular way [9], and moreover a large class of axiomatic extensions of the basic LE-logics can be endowed with analytic calculi by algorithmically generating the analytic structural rules corresponding to the additional axioms [13, 14]. This proof-theoretic environment makes it possible to study interpolation through purely proof-theoretic criteria on structural rules, in a way analogous to the role played by frame rules in Kuznets’ approach to labelled calculi.

The main contribution of the present paper is a general constructive proof of Lyndon interpolation for basic LE-logics using display calculi, together with a modular extension method for axiomatic extensions. More precisely, we show that the display calculus **D.LE** for a basic LE-logic in any given but arbitrary signature enjoys a *local Maehara interpolation property*, from which effective interpolation follows. We then identify a class of structural rules, which we call *interpolation-safe rules*, preserving this property. These rules play a role analogous to the Horn-like geometric rules in Kuznets’ framework, but are formulated intrinsically in the setting of display calculi. As a consequence, interpolation for many axiomatic extensions of LE-logics can be established uniformly by checking local syntactic conditions on the added structural rules.

Our framework applies in particular to several recently introduced lattice-based modal logics. As illustrations, we obtain constructive Lyndon interpolation results for fundamental logic and its tense modal expansions introduced in [11, 12]. These logics arise from the study of non-distributive lattice semantics motivated both by proof theory and by applications to modal reasoning and natural language semantics.

Structure of the paper. In Section 2, we recall the necessary background on LE-logics and display calculi. In Section 3, we provide the formal definitions of the various interpolation properties and establish the local Lyndon/Maehara interpolation property for the basic calculus **D.LE**. In Section 4, we show how to extend the strategy to some axiomatic extensions of our base logic and showcase some examples. We conclude in Section 5 with some directions for further work.

2. Preliminaries

LE-logics are the logics canonically associated with varieties of lattice expansions, i.e. *bounded lattices* expanded with operations that either preserve or reverse finite (hence also empty) lattice operations in each coordinate [9]. Examples are intuitionistic logic [15], bi-intuitionistic logic [16], distributive modal logic [17], positive modal logic [18], semi-De Morgan logic [19], Orthologic [20], the full Lambek calculus [21, 22], the Lambek-Grishin calculus [23], or MALL, the exponential-free fragment of linear logic [24].

The lack of distributivity is a key feature in several important applications from logics underlying quantum mechanics [25, 26] to logics of formal concepts and categories [27]. In [8], the relational semantics of (some) LE-logics are exploited to provide an intuitive explanation of why, or under which circumstances, the failure of distributivity is a reasonable and desirable feature. In [8], it is argued that LE-logics can be taken as the logics of categories and evidential reasoning, and in [28] as the logics of informational entropy.

2.1. LE-Logics and Their Algebraic Semantics

An *order-type* over $n \in \mathbb{N}$ is an n -tuple $\varepsilon \in \{1, \partial\}^n$. For every order type ε , we denote its *opposite* order type by ε^∂ , that is, $\varepsilon_i^\partial = 1$ iff $\varepsilon_i = \partial$ for every $1 \leq i \leq n$.

The language $\mathcal{L}_{\text{LE}}(\mathcal{F}, \mathcal{G})$ (from now on abbreviated as $\mathcal{L}_{\text{LE}}$) takes as parameters a denumerable set of proposition letters AtProp , elements of which are denoted p, q, r , possibly with indexes, and disjoint sets of connectives $\mathcal{F}$ and $\mathcal{G}$. Each $f \in \mathcal{F}$ (resp. $g \in \mathcal{G}$) has arity $n_f \in \mathbb{N}$ (resp. $n_g \in \mathbb{N}$) and is associated with some order-type ε_f over n_f (resp. ε_g over n_g). The terms (formulas) of $\mathcal{L}_{\text{LE}}$ are defined recursively as follows:

$$\varphi := p \mid \perp \mid \top \mid \varphi \wedge \varphi \mid \varphi \vee \varphi \mid f(\varphi_1, \dots, \varphi_{n_f}) \mid g(\varphi_1, \dots, \varphi_{n_g})$$

where $p \in \text{AtProp}$, $f \in \mathcal{F}$, $g \in \mathcal{G}$. Terms in $\mathcal{L}_{\text{LE}}$ will be denoted either by s, t , or by lowercase Greek letters such as φ, ψ, γ etc. In the remainder of the paper, when it is clear from the context, we will often simplify notation and write e.g. n for n_f and ε_i for $\varepsilon_{f,i}$. We also extend the $\{1, \partial\}$ -notation to the symbols $\vee, \wedge, \perp, \top, \leq, \vdash$ by defining

$$\vee^\partial = \wedge, \quad \wedge^\partial = \vee, \quad \perp^\partial = \top, \quad \top^\partial = \perp, \quad \leq^\partial = \geq, \quad \varphi \vdash^\partial \psi = \psi \vdash \varphi$$

while superscript 1 denotes the identity map. Therefore, in what follows, we will sometimes write e.g. $\vee^{\varepsilon_i}$ to denote $\vee$ when $\varepsilon_i = 1$ and $\wedge$ when $\varepsilon_i = \partial$. In what follows, for every $k \in \mathcal{F} \cup \mathcal{G}$ we use $k(\bar{\psi})[\varphi]_i$ to indicate that the formula φ occurs in the i -th coordinate of the vector $\bar{\psi}$.

For any $\mathcal{L}_{\text{LE}} = \mathcal{L}_{\text{LE}}(\mathcal{F}, \mathcal{G})$, the *basic*, or *minimal* $\mathcal{L}_{\text{LE}}$ -logic $\mathbf{L}_{\text{LE}}$ is a set of sequents $\varphi \vdash \psi$, with $\varphi, \psi \in \mathcal{L}_{\text{LE}}$, which contains the following sequents as axioms:

$$\begin{aligned} &\perp \vdash p, \quad p \vdash p, \quad p \vdash \top, \quad p \vdash p \vee q, \quad q \vdash p \vee q, \quad p \wedge q \vdash p, \quad p \wedge q \vdash q, \\ &f(\bar{p})[\perp^{\varepsilon_i}]_i \vdash \perp, \quad f(\bar{p})[q \vee^{\varepsilon_i} r]_i \vdash f(\bar{p})[q]_i \vee f(\bar{p})[r]_i, \\ &\top \vdash g(\bar{p})[\top^{\varepsilon_i}]_i, \quad g(\bar{p})[q]_i \wedge g(\bar{p})[r]_i \vdash g(\bar{p})[q \wedge^{\varepsilon_i} r]_i, \end{aligned}$$

and is closed under the following inference rules (recall that $\varphi \vdash^\partial \psi$ means $\psi \vdash \varphi$):

$$\begin{array}{c} \frac{\varphi \vdash \chi \quad \chi \vdash \psi}{\varphi \vdash \psi} \quad \frac{\varphi \vdash \psi}{\varphi(\chi/p) \vdash \psi(\chi/p)} \quad \frac{\chi \vdash \varphi \quad \chi \vdash \psi}{\chi \vdash \varphi \wedge \psi} \quad \frac{\varphi \vdash \chi \quad \psi \vdash \chi}{\varphi \vee \psi \vdash \chi} \\[10pt] \frac{\varphi \vdash^{\varepsilon_{f,i}} \psi}{f(\bar{p})[\varphi]_i \vdash f(\bar{p})[\psi]_i} \quad \frac{\varphi \vdash^{\varepsilon_{g,i}} \psi}{g(\bar{p})[\varphi]_i \vdash g(\bar{p})[\psi]_i} \end{array}$$

In a basic language $\mathcal{L}_{\text{LE}}(\mathcal{F}, \mathcal{G})$, the elements of $\mathcal{F} \cup \mathcal{G}$ are usually mutually independent. However, in some cases, some connectives in $\mathcal{F} \cup \mathcal{G}$ are one another's residuals (or (dual) Galois-residuals) in some coordinates.¹ In particular, given $f \in \mathcal{F}$ or $g \in \mathcal{G}$ we have $f_i^\# \in \mathcal{G}$ if $\varepsilon_{f,i} = 1$ or $g_i^\# \in \mathcal{F}$ if $\varepsilon_{g,i} = 1$, and $f_i^\# \in \mathcal{F}$ if $\varepsilon_{f,i} = \partial$ or $g_i^\# \in \mathcal{G}$ if $\varepsilon_{g,i} = \partial$, the order-type of which are (i) $\varepsilon_{f_i^\#,i} = \varepsilon_{f,i}$ and $\varepsilon_{f_i^\#,j} = (\varepsilon_{f,j})^{\varepsilon_{f,i}^\partial}$ for any $j \neq i$, and (ii) $\varepsilon_{g_i^\#,i} = \varepsilon_{g,i}$ and $\varepsilon_{g_i^\#,j} = (\varepsilon_{g,j})^{\varepsilon_{g,i}^\partial}$ for any $j \neq i$. For instance, if

¹In what follows, we will typically not make a difference between co-variant and contravariant residuation, and refer to any of these as 'residuals'.

f and g are binary connectives such that $\varepsilon_f = (1, \partial)$ and $\varepsilon_g = (\partial, 1)$, then $\varepsilon_{f_1^\#} = (1, 1)$, $\varepsilon_{f_2^\#} = (1, \partial)$, $\varepsilon_{g_1^\#} = (\partial, 1)$ and $\varepsilon_{g_2^\#} = (1, 1)$.² The intended interpretation of the n_f -ary connective $f_i^\#$ is the right residual of $f \in \mathcal{F}$ in its i th coordinate if $\varepsilon_{f,i} = 1$ (resp. its dual Galois-residual if $\varepsilon_{f,i} = \partial$). The intended interpretation of the n_g -ary connective $g_i^\#$ is the left residual of $g \in \mathcal{G}$ in its i th coordinate if $\varepsilon_{g,i} = 1$ (resp. its Galois-residual if $\varepsilon_{g,i} = \partial$). If so, $\mathbf{L}_{LE}$ is augmented with the invertible rules

$$\frac{f(\bar{\phi})[\varphi]_i \vdash \psi}{\varphi \vdash^{\varepsilon_{f,i}} f_i^\#(\bar{\phi})[\psi]_i} \quad \frac{\varphi \vdash g(\bar{\phi})[\psi]_i}{g_i^\#(\bar{\phi})[\varphi]_i \vdash^{\varepsilon_{g,i}} \psi}$$

Any language $\mathcal{L}_{LE} = \mathcal{L}_{LE}(\mathcal{F}, \mathcal{G})$ can be associated with the language $\mathcal{L}_{LE}^* = \mathcal{L}_{LE}(\mathcal{F}^*, \mathcal{G}^*)$, where $\mathcal{F}^* \supseteq \mathcal{F}$ and $\mathcal{G}^* \supseteq \mathcal{G}$ are obtained by expanding $\mathcal{L}_{LE}$ with the residuals of each connective in each coordinate. Then, the logic $\mathbf{L}_{LE}$ can be expanded to the minimal fully residuated $\mathcal{L}_{LE}^*$ -logic $\mathbf{L}_{LE}^*$, by adding the corresponding residuation rules. The logic $\mathbf{L}_{LE}^*$ is a conservative extension of $\mathbf{L}_{LE}$ (cf. [14, Theorem 2.4]), i.e. every $\mathcal{L}_{LE}$ -sequent $\varphi \vdash \psi$ is derivable in $\mathbf{L}_{LE}$ iff $\varphi \vdash \psi$ is derivable in $\mathbf{L}_{LE}^*$.

Example 2.1. *Fundamental modal logic [12] is an LE-logic in the signature $\mathcal{F} := \{\Diamond\}$, $\mathcal{G} := \{\neg, \Box\}$, where $n_\neg = n_\Box = n_\Diamond = 1$, $\varepsilon_\Box = \varepsilon_\Diamond = 1$ and $\varepsilon_\neg = \partial$. It is obtained by adding the following sequents and rules to $\mathbf{L}_{LE}$.*

$$p \wedge \neg p \vdash \perp \quad \frac{\varphi \vdash \neg \psi}{\psi \vdash \neg \varphi} \quad \Diamond \neg p \vdash \neg \Box p$$

In this case, $\mathcal{F}^* = \{\Diamond, \blacklozenge\}$ and $\mathcal{G}^* = \{\neg, \Box, \blacksquare\}$, where $\blacklozenge$ is the left residual of $\Box$ and $\blacksquare$ is the right residual of $\Diamond$. Notice that $\neg$ is its own Galois-residual. In the following, we write tense fundamental modal logic to denote Holliday's fundamental modal logic augmented with the residuals of the modal connectives.

2.2. Proper Display Calculi for Basic Normal LE-Logics

In this section, we recall the definition of the proper display calculus $\mathbf{D.LE}$ for the basic normal $\mathcal{L}_{LE}$ -logic for a fixed but arbitrary LE-signature $\mathcal{L} = \mathcal{L}_{LE}(\mathcal{F}, \mathcal{G})$ (cf. Section 2.1). Our presentation is a more streamlined version of the one introduced in [13] for DLE-logics and then straightforwardly generalized to LE-logics in [14].

Let $S_{\mathcal{F}} := \{\hat{f} \mid f \in \mathcal{F}^*\}$ and $S_{\mathcal{G}} := \{\check{g} \mid g \in \mathcal{G}^*\}$ be the sets of structural connectives indexed by $\mathcal{F}^*$ and $\mathcal{G}^*$ respectively. Each structural connective comes with an arity and an order-type which coincides with those of its associated operational connective in $\mathcal{F}^*$ and $\mathcal{G}^*$. The calculus $\mathbf{D.LE}$ manipulates sequents $\Pi \vdash \Sigma$, where Π and Σ are structures of two sorts which are built from formulas and are defined by the following simultaneous recursion:

$$\text{Str}_{\mathcal{F}} \ni \Pi := A \mid \hat{\top} \mid \hat{f}(\bar{\Xi}) \quad \text{Str}_{\mathcal{G}} \ni \Sigma := A \mid \check{\perp} \mid \check{g}(\bar{\Xi})$$

where A is an $\mathcal{L}_{LE}$ -formula, $\hat{f} \in S_{\mathcal{F}}$, $\check{g} \in S_{\mathcal{G}}$, and $\bar{\Xi} \in \text{Str}_{\mathcal{F}}^{\varepsilon_f}$ (resp. $\bar{\Xi} \in \text{Str}_{\mathcal{G}}^{\varepsilon_g}$), where $\text{Str}_{\mathcal{F}}^{\varepsilon_f} = \text{Str}_{\mathcal{F}}^{\varepsilon_{f,1}} \times \dots \times \text{Str}_{\mathcal{F}}^{\varepsilon_{f,n_f}}$ where $\text{Str}_{\mathcal{F}}^{\partial} = \text{Str}_{\mathcal{G}}$, and dually for $\text{Str}_{\mathcal{G}}^{\varepsilon_g}$. In what follows, for every $o \in S_{\mathcal{F}} \cup S_{\mathcal{G}}$ we use $o(\bar{\Xi})[\Pi]_i$ (resp. $o(\bar{\Xi})[\Sigma]_i$) to indicate that the structure Π (resp. Σ) occurs in the i -th coordinate of the vector $\bar{\Xi}$.

If S_1 and S_2 are sequents, the double-horizontal-line notation $r \frac{S_1}{S_2}$ is an abbreviation for the rules $r \frac{S_1}{S_2}$ and $\frac{S_2}{S_1} r^{-1}$.

- Identity and Cut rules:

²Note that this notation depends on the connective which is taken as primitive, and needs to be carefully adapted to well known cases. For instance, consider the 'fusion' connective $\circ$ (which, when denoted as f , is such that $\varepsilon_f = (1, 1)$). Its residuals $f_1^\#$ and $f_2^\#$ are commonly denoted $/$ and $\backslash$ respectively. However, if $\backslash$ is taken as the primitive connective g , then $g_2^\#$ is $\circ = f$, and $g_1^\#(x_1, x_2) := x_2/x_1 = f_1^\#(x_2, x_1)$. This example shows that, when identifying $g_1^\#$ and $f_1^\#$, the conventional order of the coordinates is not preserved, and depends on which connective is taken as primitive.

$$\text{Id}_p \frac{}{p \vdash p} \quad \frac{\Pi \vdash A \quad A \vdash \Sigma}{\Pi \vdash \Sigma} \text{Cut}_A$$

- Display rules for $f \in \mathcal{F}$ and $g \in \mathcal{G}$: for any $1 \leq k \leq n_f$ and $1 \leq \ell \leq n_g$,³

If $\varepsilon_{f,k} = 1$ and $\varepsilon_{g,\ell} = 1$

$$\hat{f} \dashv \check{f}_k^\# \frac{\hat{f}(\bar{\Xi})[\Pi]_k \vdash \Sigma}{\Pi \vdash \check{f}_k^\#(\bar{\Xi})[\Sigma]_k} \quad \frac{\Pi \vdash \check{g}(\bar{\Xi})[\Sigma]_\ell}{\hat{g}_\ell^b(\bar{\Xi})[\Pi]_\ell \vdash \Sigma} \hat{g}_\ell^b \dashv \check{g}$$

If $\varepsilon_{f,k} = \partial$ and $\varepsilon_{g,\ell} = \partial$

$$(\hat{f}, \hat{f}_k^\#) \frac{\hat{f}(\bar{\Xi})[\Sigma]_k \vdash \Sigma'}{\hat{f}_k^\#(\bar{\Xi})[\Sigma']_k \vdash \Sigma} \quad \frac{\Pi' \vdash \check{g}(\bar{\Xi})[\Pi]_\ell}{\Pi \vdash \check{g}_j^b(\bar{\Xi})[\Pi']_\ell} (\check{g}, \check{g}_\ell^b)$$

- Logical introduction rules for lattice connectives, for $i \in \{1, 2\}$:

$$\begin{array}{c} \top_L \frac{\top \vdash \Sigma}{\top \vdash \Sigma} \quad \frac{}{\Pi \vdash \top} \top_R \quad \wedge_{Li} \frac{A_i \vdash \Sigma}{A_1 \wedge A_2 \vdash \Sigma} \quad \frac{\Pi \vdash A \quad \Pi \vdash B}{\Pi \vdash A \wedge B} \wedge_R \\ \\ \perp_L \frac{}{\perp \vdash \Sigma} \quad \frac{\Pi \vdash \perp}{\Pi \vdash \perp} \perp_R \quad \vee_L \frac{A \vdash \Sigma \quad B \vdash \Sigma}{A \vee B \vdash \Sigma} \quad \frac{\Pi \vdash A_i}{\Pi \vdash A_1 \vee A_2} \vee_{Ri} \end{array}$$

- Logical introduction rules for $f \in \mathcal{F}$ and $g \in \mathcal{G}$:

$$f_L \frac{\hat{f}(\bar{A}) \vdash \Sigma}{f(\bar{A}) \vdash \Sigma} \quad \frac{(\Xi_i \vdash^{\varepsilon_{f,i}} A_i)_{1 \leq i \leq n_f}}{\hat{f}(\bar{\Xi}) \vdash f(\bar{A})} f_R \quad g_L \frac{(A_i \vdash^{\varepsilon_{g,i}} \Xi_i)_{1 \leq i \leq n_g}}{g(\bar{A}) \vdash \check{g}(\bar{\Xi})} \quad \frac{\Pi \vdash \check{g}(\bar{A})}{\Pi \vdash g(\bar{A})} g_R$$

where $\bar{\Xi} = (\Xi_1, \dots, \Xi_{n_f})$ in f_R and $\bar{\Xi} = (\Xi_1, \dots, \Xi_{n_g})$ in g_L . In particular, if f and g are 0-ary (i.e. they are constants), the rules f_R and g_L above reduce to the axioms (0-ary rules) $\hat{f} \vdash f$ and $g \vdash \check{g}$.

The calculus **D.LE** is sound w.r.t. the class of complete $\mathcal{L}$ -algebras (cf. [13, Theorem 12] and [9, Theorem 2.8]). Moreover, **D.LE** is a proper display calculus (cf. [13, Theorem 26] and [9, Appendix A]), and hence cut elimination holds for it as a consequence of a Belnap-style cut elimination metatheorem (cf. [13, Section 2.2 and Appendix A] and [9, Theorem A.7]).

Example 2.2 (Continuing Example 2.1). *In the language of fundamental logic discussed in Example 2.1, the display rules are as follows:*

$$\hat{\diamond} \dashv \blacksquare \frac{\hat{\diamond}\Pi \vdash \Sigma}{\Pi \vdash \blacksquare\Sigma} \quad \hat{\diamond} \dashv \check{\square} \frac{\hat{\diamond}\Pi \vdash \Sigma}{\Pi \vdash \check{\square}\Sigma} \quad (\sim, \sim) \frac{\Pi_1 \vdash \sim\Pi_2}{\Pi_2 \vdash \sim\Pi_1}$$

Notice that the unary rule of fundamental logic presented in Example 2.1 corresponds to the third display rule above. The logical introduction rules for the modalities and the negation of fundamental logic are:

$$\diamond_L \frac{\hat{\diamond}A \vdash \Sigma}{\diamond A \vdash \Sigma} \quad \frac{\Pi \vdash A}{\hat{\diamond}\Pi \vdash \diamond A} \diamond_R \quad \square_L \frac{A \vdash \Sigma}{\square A \vdash \check{\square}\Sigma} \quad \frac{\Pi \vdash \check{\square}A}{\Pi \vdash \square A} \square_R \quad \neg_L \frac{\Pi \vdash A}{\neg A \vdash \sim\Pi} \quad \frac{\Pi \vdash \sim A}{\Pi \vdash \neg A} \neg_R$$

Finally, in the case of tense fundamental logic, the following introduction rules are added:

$$\diamond_L \frac{\hat{\diamond}A \vdash \Sigma}{\diamond A \vdash \Sigma} \quad \frac{\Pi \vdash A}{\hat{\diamond}\Pi \vdash \diamond A} \diamond_R \quad \blacksquare_L \frac{A \vdash \Sigma}{\blacksquare A \vdash \blacksquare\Sigma} \quad \frac{\Pi \vdash \blacksquare A}{\Pi \vdash \blacksquare A} \blacksquare_R$$

³The notation $\hat{f} \dashv \check{f}_k^\#$ (resp. $\hat{g}_\ell^b \dashv \check{g}$) indicates that $\hat{f}$ and $\check{f}_k^\#$ (resp. $\hat{g}_\ell^b$ and $\check{g}$) are in a *residuated pair* and $\check{f}_k^\#$ (resp. $\hat{g}_\ell^b$) is the right residual (resp. left residual) of $\hat{f}$ (resp. $\check{g}$) in the k -th coordinate (resp. ℓ -th coordinate). The notation $(\check{g}, \check{g}_\ell^b)$ (resp. $(\hat{f}, \hat{f}_k^\#)$) indicates that $\check{g}$ and $\check{g}_\ell^b$ (resp. $\hat{f}$ and $\hat{f}_k^\#$) are in a *Galois connection* (resp. *dual Galois connection*) and $\check{g}_\ell^b$ (resp. $\hat{f}_k^\#$) is the right residual (resp. left residual) of $\check{g}$ (resp. $\hat{f}$) in the ℓ -th coordinate (resp. k -th coordinate).

Metastructures and Structural Rules. In the presentation of the language of **D.LE**, Σ_i , Π_i , and Ξ_i are metavariables for structures. To formally present analytic structural rules, we need to make use of metastructures, i.e., structures that are constructed by structural metavariables. Let $\text{MVar} = \text{MVar}_{\mathcal{F}} \uplus \text{MVar}_{\mathcal{G}}$ be the denumerable set of *metavariables* of sorts $\Pi_1, \Pi_2, \dots \in \text{MVar}_{\mathcal{F}}$ and $\Sigma_1, \Sigma_2, \dots \in \text{MVar}_{\mathcal{G}}$. The sets $\text{MStr}_{\mathcal{F}}$ and $\text{MStr}_{\mathcal{G}}$ of the $\mathcal{F}$ - and $\mathcal{G}$ -metastructures are defined by simultaneous induction as follows:

$$\text{MStr}_{\mathcal{F}} \ni \Phi := \Pi \mid \hat{f}(\bar{\Phi}) \quad \text{MStr}_{\mathcal{G}} \ni \Psi := \Sigma \mid \check{g}(\bar{\Psi})$$

where $\Pi \in \text{MVar}_{\mathcal{F}}$, $\Sigma \in \text{MVar}_{\mathcal{G}}$, $\hat{f} \in \mathcal{F}^*$ and $\check{g} \in \mathcal{G}^*$ and $\bar{\Phi} \in \text{MStr}_{\mathcal{F}}^{\varepsilon_f}$, and $\bar{\Psi} \in \text{MStr}_{\mathcal{G}}^{\varepsilon_g}$, and for any order-type ε on n , we let $\text{MStr}_{\mathcal{F}}^{\varepsilon} := \prod_{i=1}^n \text{MStr}_{\mathcal{F}}^{\varepsilon_i}$ and $\text{MStr}_{\mathcal{G}}^{\varepsilon} := \prod_{i=1}^n \text{MStr}_{\mathcal{G}}^{\varepsilon_i}$, where for all $1 \leq i \leq n$,

$$\text{MStr}_{\mathcal{F}}^{\varepsilon_i} = \begin{cases} \text{MStr}_{\mathcal{F}} & \text{if } \varepsilon_i = 1 \\ \text{MStr}_{\mathcal{G}} & \text{if } \varepsilon_i = \partial \end{cases} \quad \text{MStr}_{\mathcal{G}}^{\varepsilon_i} = \begin{cases} \text{MStr}_{\mathcal{G}} & \text{if } \varepsilon_i = 1, \\ \text{MStr}_{\mathcal{F}} & \text{if } \varepsilon_i = \partial. \end{cases}$$

Definition 2.3. An analytic structural rule is a rule of the form

$$\frac{(\Phi_1 \vdash \Psi_1)(\bar{\Upsilon}_1) \quad \dots \quad (\Phi_n \vdash \Psi_n)(\bar{\Upsilon}_n)}{(\Phi_0 \vdash \Psi_0)(\bar{\Upsilon}_0)}$$

where $\bar{\Upsilon}_i$ with $0 \leq i \leq n$ is the set of metavariables occurring in each sequent $\Phi_i \vdash \Psi_i$, $\bar{\Upsilon}_0 \supseteq \bar{\Upsilon}_1 \cup \dots \cup \bar{\Upsilon}_n$, and while structural metavariables might occur multiple times in the premises they occur only once in the conclusion. An instance of the rule R is obtained from R by uniformly substituting each structural metavariable $\Pi \in \text{MVar}_{\mathcal{F}}$ with an element of $\text{Str}_{\mathcal{F}}$ and every structural metavariable $\Sigma \in \text{MVar}_{\mathcal{G}}$ with an element of $\text{Str}_{\mathcal{G}}$. A calculus contains the analytic rule R if it contains every instance of R .

The expressivity of proper display calculi has been characterized in terms of the *analytic inductive* axioms (cf. [13, Definition 55]), a proper subclass of the class of axioms to which the state of the art algorithmic correspondence theory via second-order quantifier elimination applies:

Proposition 2.4. (cf. [13, Proposition 59, Proposition 61]) Every analytic inductive LE-inequality can be equivalently transformed, via an ALBA-reduction, into a set of analytic structural rules. Conversely, every analytical structural rule is equivalent to an inductive LE-inequality.

Example 2.5 (Continuing Examples 2.1, 2.2). The axioms $p \wedge \neg p \vdash \perp$ and $\Diamond \neg p \vdash \neg \Box p$ correspond, respectively, to the following two rules

$$R_1 \frac{X \vdash \neg X}{X \vdash Y} \quad R_2 \frac{X \vdash \neg \Diamond Y}{X \vdash \Box \neg Y}$$

2.3. Auxiliary Notations and Lemmas

In this section, we gather some notation and facts that will be useful throughout the paper.

Definition 2.6. The **D.LE**-sequents $\Pi_1 \vdash \Sigma_1$ and $\Pi_2 \vdash \Sigma_2$ are display equivalent (notation: $\Pi_1 \vdash \Sigma_1 \equiv_D \Pi_2 \vdash \Sigma_2$) if $\Pi_1 \vdash \Sigma_1$ can be derived from $\Pi_2 \vdash \Sigma_2$ using only display postulates.

It immediately follows from the definition that $\equiv_D$ is an equivalence relation.

When a structure is denoted with a capital letter, e.g. $X, U, \Pi, \dots$, we will denote with the lower-case counterpart $x, u, \pi, \dots$ the formula that arises by substituting the structural connectives of the structure with their operational counterparts.

We will call a structure X a substructure of the sequent $\Pi \vdash \Sigma$ if X is a substructure of Π or a substructure of Σ , and denote this relationship by $(\Pi \vdash \Sigma)[X]$. In this notation, X is understood to always refer to a concrete instance of the substructure X in the sequent $\Pi \vdash \Sigma$. We will write $(\Pi \vdash \Sigma)[Y/X]$ to denote the substitution of the instance of the substructure X in $\Pi \vdash \Sigma$ with the structure Y .

The *display property* states that for every substructure X of $\Pi \vdash \Sigma$, there exists a display equivalent sequent either of the form $X \vdash \Sigma'$ or of the form $\Pi' \vdash X$, such that for any $\mathcal{F}$ -structure Γ (resp. $\mathcal{G}$ -structure Δ) the sequent $\Gamma \vdash \Sigma'$ (resp. $\Pi' \vdash \Delta$) is display equivalent to $(\Pi \vdash \Sigma)[\Gamma/X]$ (resp. $(\Pi \vdash \Sigma)[\Delta/X]$).

Given $(\Pi \vdash \Sigma)[X]$, we define $\varepsilon_X^{\Pi \vdash \Sigma} = 1$ if $\Pi \vdash \Sigma$ is display equivalent to $X \vdash \Sigma'$ and $\varepsilon_X^{\Pi \vdash \Sigma} = \partial$ if $\Pi \vdash \Sigma$ is display equivalent to $\Pi' \vdash X$. In this notation, we will often omit the superscript when it is clear from the context.

The following lemma shows the inverse direction of the display property:

Lemma 2.7. *If $\Pi \vdash \Sigma$ and $\Pi' \vdash \Sigma'$ are display-equivalent, then one of the following cases holds.*

1. *The structure Π' is a substructure of $\Pi \vdash \Sigma$ and for every $\mathcal{F}$ -structure Γ , $\Gamma \vdash \Sigma'$ is display equivalent to $\Pi \vdash \Sigma[\Gamma/\Pi']$.*
2. *The structure Σ' is a substructure of $\Pi \vdash \Sigma$ and for every $\mathcal{G}$ -structure Δ , $\Pi' \vdash \Delta$ is display equivalent to $\Pi \vdash \Sigma[\Delta/\Sigma']$.*

Proof. By assumption, a derivation exists of $\Pi' \vdash \Sigma'$ from $\Pi \vdash \Sigma$ that consists only of applications of display postulates. We proceed by induction on the length n of this derivation with the following stronger induction hypothesis: In $\Pi' \vdash \Sigma'$ the instantiations of all the metavariables with the exception of exactly one that appear in the last display rule that was applied, are substructures of $\Pi \vdash \Sigma$. If X_0 is the metavariable whose instantiation in $\Pi' \vdash \Sigma'$, Z_0 , is not a substructure of $\Pi \vdash \Sigma$, then, by applying to $\Pi' \vdash \Sigma'$ the display rule that puts X_0 on display, we obtain a sequent $Z_0 \vdash^{\varepsilon_{X_0}} \Delta$, that has already appeared in the proof as a sequent strictly before the last one.

Before moving to the proof by induction, let us first show that if the hypothesis is true, then the desideratum follows.

We can assume without loss of generality that every sequent appears only once in the proof; otherwise, the proof can be shortened. In this context in particular, the induction hypothesis implies that the instantiation of the structure meta-variables in display after the application of a display rule, throughout the proof, and so also in the last application of the rule, is a substructure of $\Pi \vdash \Sigma$. This, along with the induction hypothesis, implies that the structure that instantiates the meta-variable on display in the last application of a display rule, in each step of the proof is a substructure of a structure that instantiates a meta-variable of the applied rule. This is because, otherwise it must contain a substructure that is not a substructure of $\Pi \vdash \Sigma$, contradicting the fact that it's a substructure of $\Pi \vdash \Sigma$. Now, let us assume that X is the structure that instantiates the metavariable in display in the last rule application, and Z a structure of the same type ($\mathcal{F}$ or $\mathcal{G}$). Substituting in each sequent in the proof Z for X , does not affect the validity of the proof, since X in each step is a substructure of a structure that instantiates a metavariable. As display rules are bidirectional, we obtain the desideratum.

Now we can proceed with the induction proof. The statement holds vacuously for a proof of size 0. Let us assume that $n \geq 1$, and that the statement holds for all derivations of length $k < n$. If the last rule applied is $\hat{f} \dashv \check{f}_i^\sharp$, then the derivation looks as follows:

$$\frac{\frac{\frac{\Pi \vdash \Sigma}{\vdots}}{\hat{f}(\bar{\Xi})[\Pi']_i \vdash T}}{\Pi' \vdash \check{f}_i^\sharp(\bar{\Xi})[T]_i}$$

If $\hat{f}(\bar{\Xi})[\Pi']_i$ corresponded to an instantiation of a metavariable in the previous application of a display rule, then by the induction hypothesis, it is a substructure of $\Pi \vdash \Sigma$, in which case $\bar{\Xi}$ are all substructures of $\Pi \vdash \Sigma$, and so is Π' and applying a display rule that puts T on display yields a sequent that has already appeared in the proof, namely $\hat{f}(\bar{\Xi})[\Pi']_i \vdash T$.

If $\hat{f}(\bar{\Xi})[\Pi']_i$ doesn't correspond to an instantiation of a metavariable in the previous application of a display rule, then Π' , $\bar{\Xi}$, and T do. In this case, the instantiation that doesn't correspond to a substructure of $\Pi \vdash \Sigma$ cannot be Π' , since then the sequent $\Pi' \vdash \check{f}_i^\sharp(\bar{\Xi})[T]_i$ has already appeared in the

proof, a possibility we have excluded. Hence, the required property also holds in this case. In case of different display rules we argue similarly. This concludes the induction step and the proof. $\square$

3. Interpolation Properties in Display Calculi

In this section, we introduce the refinements of the Lyndon and Maehara [29, 30] interpolation notions in the context of display calculi, and prove that in this context the two notions coincide. We show that the rules of the basic normal LE-logic $\mathbf{L}_{LE}$ in any LE-language preserve the existence of such interpolants, made precise through the notion of the local interpolation property. Through these results, we show that the Maehara interpolation properties hold for $\mathbf{L}_{LE}$ in any LE-language.

3.1. Lyndon and Maehara Interpolation

Definition 3.1. A Lyndon interpolant for a $\mathbf{D.LE}$ sequent $\Pi \vdash \Sigma$ is an LE-formula γ such that both $\Pi \vdash \gamma$ and $\gamma \vdash \Sigma$ are derivable in $\mathbf{D.LE}$; moreover, $\text{Var}^+(\gamma) \subseteq \text{Var}^+(\Pi) \cap \text{Var}^+(\Sigma)$, and $\text{Var}^-(\gamma) \subseteq \text{Var}^-(\Pi) \cap \text{Var}^-(\Sigma)$, where $\text{Var}^+(\Upsilon)$ (resp. $\text{Var}^-(\Upsilon)$) is the set of the atomic variables occurring positively (resp. negatively) in Υ . The Lyndon interpolation property holds for a sequent $\Pi \vdash \Sigma$ if a Lyndon interpolant for it exists.

Definition 3.2. The Maehara interpolation property holds for a $\mathbf{D.LE}$ sequent $\Pi \vdash \Sigma$ if for every substructure X occurring in the sequent, some $\mathcal{L}$ -formula γ_X exists such that both $X \vdash^{\varepsilon_X} \gamma_X$ and $(\Pi \vdash \Sigma)[\gamma_X/X]$ are derivable in $\mathbf{D.LE}$, and $\text{Var}^+(\gamma_X) \subseteq \text{Var}^+(X) \cap \text{Var}^+(\Pi \vdash \Sigma)[[-]/X]$, and $\text{Var}^-(\gamma_X) \subseteq \text{Var}^-(X) \cap \text{Var}^-(\Pi \vdash \Sigma)[[-]/X]$.

Definition 3.3. The local Lyndon/Maehara interpolation property holds for a calculus if it is preserved by each rule of the calculus; i.e., if it holds for (any instantiation of) the premises of any given rule, then it does for (the instantiation of) the conclusion.

Lemma 3.4. A $\mathbf{D.LE}$ sequent $\Pi \vdash \Sigma$ has the Maehara interpolation property if and only if every sequent, $\Pi' \vdash \Sigma'$, display equivalent to $\Pi \vdash \Sigma$, has a Lyndon interpolant.

Proof. Assume that $\Pi \vdash \Sigma$ has the local Maehara interpolation property and let $\Pi' \vdash \Sigma'$ be display equivalent to $\Pi \vdash \Sigma$. By Lemma 2.7, it follows that either Π' or Σ' is a substructure of $\Pi \vdash \Sigma$. Let's assume that we are in case 1 of Lemma 2.7, and so, in particular, Π' is a substructure of $\Pi \vdash \Sigma$. By assumption, there exists an interpolant γ such that $\Pi' \vdash \gamma$ and $(\Pi \vdash \Sigma)[\gamma/\Pi']$ are derivable. Then by Lemma 2.7 we have that $\gamma \vdash \Sigma'$ is display equivalent to $(\Pi \vdash \Sigma)[\gamma/\Pi']$. Hence, γ is a Lyndon interpolant of $\Pi' \vdash \Sigma'$. For the other case we work analogously.

Now assume that every sequent, $\Pi' \vdash \Sigma'$, display equivalent to $\Pi \vdash \Sigma$ has a Lyndon interpolant, and let X be a substructure of $\Pi \vdash \Sigma$. Let's assume without loss of generality that $\varepsilon_\Delta = 1$. By the display property, there exists a sequent, $X \vdash \Sigma'$, which is display equivalent to $\Pi \vdash \Sigma$, such that for every $\mathcal{F}$ -substructure Γ , $\Gamma \vdash \Sigma'$ is display equivalent to $(\Pi \vdash \Sigma)[\Gamma/X]$. By assumption, $X \vdash \Sigma'$ has a Lyndon interpolant γ , hence $X \vdash \gamma$ and $\gamma \vdash \Sigma'$ are derivable. The second sequent is display equivalent to $(\Pi \vdash \Sigma)[\gamma/X]$. Hence, $\Pi \vdash \Sigma$ has the Maehara interpolation property. $\square$

3.2. The Interpolation Result

Lemma 3.5. The local Lyndon interpolation property holds for a display calculus if and only if the local Maehara interpolation property holds for it.

Proof. If the local Lyndon interpolation property holds for a display calculus D , then the display rules of D preserve the local Lyndon interpolation property. Hence, if the Lyndon interpolation property holds for a sequent $\Pi \vdash \Sigma$, it also holds for all its display-equivalent sequents, and hence, by Lemma 3.4, the Maehara interpolation property also holds for $\Pi \vdash \Sigma$. Hence, the Maehara interpolation property is preserved by the rules of the calculus. $\square$

Lemma 3.6. *Let*

$$R \frac{\Phi_1 \vdash \Psi_1 \quad \dots \quad \Phi_n \vdash \Psi_n}{\Phi_0 \vdash \Psi_0}$$

be a rule of a display calculus such that each structural metavariable occurs at most once in each premise⁴. Consider an instantiation of R ,

$$R_{inst} \frac{\Pi_1 \vdash \Sigma_1 \quad \dots \quad \Pi_n \vdash \Sigma_n}{\Pi_0 \vdash \Sigma_0}$$

and let Γ be an $\mathcal{F}$ -substructure of the instantiation of a metavariable X of R in R_{inst} . If, for every $\Phi_i \vdash \Psi_i$ in which X occurs, some γ_i exists such that $\Gamma \vdash \gamma_i$ and $(\Pi_i \vdash \Sigma_i)[\gamma_i/\Gamma]_k$ are derivable, then some γ_0 exists such that $\Gamma \vdash \gamma_0$ and $(\Pi_0 \vdash \Sigma_0)[\gamma_0/\Gamma]_k$ are derivable. The dual statement holds for a $\mathcal{G}$ -substructure Δ .

Proof. Let $S = \{i \mid \text{the metavariable } X \text{ occurs in } \Phi_i \vdash \Psi_i\}$. Let $\gamma_0 := \bigwedge_{i \in S} \gamma_i$. Using the right introduction for $\wedge$ we conclude that $\Gamma \vdash \bigwedge_{i \in S} \gamma_i$. Using display rules and the left introduction of $\wedge$ we obtain $(\Pi_i \vdash \Sigma_i)[\bigwedge_{i \in S} \gamma_i/\Gamma]_k$ for every $i \in S$. Applying R to these assumptions we conclude $(\Pi_0 \vdash \Sigma_0)[\gamma_0/\Gamma]_k$. If Δ is a $\mathcal{G}$ -substructure of the instantiation of a metavariable X , we use the introduction rules for $\vee$ to obtain $\delta_0 = \bigvee_{i \in S} \delta_i$. $\square$

Theorem 3.7. *The local Maehara interpolation property holds for the calculus **D.LE** associated with the basic LE-logic of any LE-language $\mathcal{L}$.*

Proof. We need to show that the local Maehara interpolation property holds for each rule of **D.LE**. Notice that all the basic rules satisfy the condition of Lemma 3.6. Hence we only need concern ourselves with structures in the conclusion that are not substructures of structural metavariables. First, let us consider the zero-ary rules:

$$\bullet \quad \text{Id} \frac{}{p \vdash p} \quad \frac{}{p \vdash \top} \top_R \quad \frac{}{\perp \vdash p} \perp_L$$

There are no proper substructures in these sequents, hence Lyndon and Maehara interpolation coincide. The local interpolants are p , $\top$, and $\perp$, respectively, which clearly also satisfy the Lyndon conditions for the variables.

Now we need to consider all the other rules and show that, assuming there is an interpolant for each premise, how to construct the interpolant for the conclusion. Let us consider the display postulates for a structural $\mathcal{F}$ -connective which is monotone in its k coordinate, i.e. $\varepsilon_{f,k} = 1$:

$$\bullet \quad \hat{f} \dashv \check{f}_k^\# \frac{\hat{f}(\Xi)[\Pi]_k \vdash \Sigma}{\Pi \vdash \check{f}_k^\#(\Xi)[\Sigma]_k}$$

The only case not covered by Lemma 3.6 is when $\Delta = \check{f}_k^\#(\Xi)[\Sigma]_k$. In this case, by assumption we have an interpolant γ such that $\Pi \vdash \gamma$ and $\hat{f}(\Xi)[\gamma]_k \vdash \Sigma$ are derivable. Then, by applying the display rule $\hat{f} \dashv \check{f}_k^\#$ from premise to conclusion, we have that $\gamma \vdash \check{f}_k^\#(\Xi)[\Sigma]_k$. Given we already know that $\Pi \vdash \gamma$, it follows that γ is also an interpolant of the sequent in the conclusion. Hence, this display rule preserves the local Maehara interpolation property. The other display rules are shown analogously.

We now consider the binary rules introducing conjunction and disjunction:

$$\bullet \quad \vee_L \frac{\varphi \vdash \Sigma \quad \psi \vdash \Sigma}{\varphi \vee \psi \vdash \Sigma} \quad \wedge_R \frac{\Pi \vdash \varphi \quad \Pi \vdash \psi}{\Pi \vdash \varphi \wedge \psi}$$

Let us focus on the right introduction $\wedge_R$. The only case not covered by Lemma 3.6 is when $\Delta = \varphi \wedge \psi$. Clearly, $\delta_\varphi \wedge \delta_\psi$ is the interpolant, where δ_φ is the interpolant for φ and δ_ψ is the interpolant for ψ in the premises. The proof for the $\vee_L$ is analogous.

⁴Notice that here we do not refer to analytic or structural rules only, hence the basic logical introduction rules are also included.

$$\bullet \quad \wedge_{L1} \frac{\varphi \vdash \Sigma}{\varphi \wedge \psi \vdash \Sigma} \quad \wedge_{L2} \frac{\psi \vdash \Sigma}{\varphi \wedge \psi \vdash \Sigma} \quad \vee_{R1} \frac{\Pi \vdash \varphi}{\Pi \vdash \varphi \vee \psi} \quad \vee_{R2} \frac{\Pi \vdash \psi}{\Pi \vdash \varphi \vee \psi}$$

For the left introduction of $\wedge$, the only non-trivial substructure is $\varphi \wedge \psi$. But then the same interpolant for ψ in the assumption works. The other rules are shown similarly.

$$\bullet \quad \top_L \frac{\hat{\top} \vdash \Sigma}{\top \vdash \Sigma} \quad \perp_R \frac{\Pi \vdash \check{\perp}}{\Pi \vdash \perp} \quad f_L \frac{\hat{f}(\bar{\varphi}) \vdash \Sigma}{f(\bar{\varphi}) \vdash \Sigma} \quad g_R \frac{\Pi \vdash \check{g}(\bar{\varphi})}{\Pi \vdash g(\bar{\varphi})}$$

These cases are similar to the previous case, and are hence omitted.

$$\bullet \quad \top_W \frac{\hat{\top} \vdash \Sigma}{\Pi \vdash \Sigma} \quad \perp_W \frac{\Pi \vdash \check{\perp}}{\Pi \vdash \Sigma}$$

The only case not covered by Lemma 3.6 in $\top_W$ is when Δ is a substructure of Π . Then the interpolant is $\top^{\varepsilon\Delta}$. Indeed, $\Delta \vdash^{\varepsilon\Delta} \top^{\varepsilon\Delta}$ is derivable, and so is $(\Pi \vdash \Sigma)[\top^{\varepsilon\Delta}/\Delta]$, by an application of $\top_W$. The case for $\perp_W$ is similar.

$$\bullet \quad \frac{\left(\Upsilon_i \vdash \varphi_i \quad \varphi_j \vdash \Upsilon_j \mid 1 \leq i, j \leq n_f, \varepsilon_f(i) = 1 \text{ and } \varepsilon_f(j) = \partial \right)}{\hat{f}(\Upsilon_1, \dots, \Upsilon_{n_f}) \vdash f(\varphi_1, \dots, \varphi_{n_f})} f_R$$

The only cases not covered by Lemma 3.6 are for the substructures $\hat{f}(\Upsilon_1, \dots, \Upsilon_{n_f})$ and $f(\varphi_1, \dots, \varphi_{n_f})$. They will have the same interpolant, in particular, if σ_i is the interpolant of Y_i , then the interpolant is $f(\sigma_1, \dots, \sigma_{n_f})$. Verifying the conditions is routine. This concludes the proof. $\square$

Corollary 3.8 (Local interpolation). *For any LE-language $\mathcal{L}_{LE}$, any D.LE-derivable sequent has the local Lyndon/Maehara interpolation property, and the interpolants can be effectively computed.*

Proof. Let $X \vdash Y$ be a derivable sequent and proceed by induction on the height of the derivation. The base cases are treated as in the proof of Theorem 3.7. For the inductive step, if R is the last rule applied in the derivation of $X \vdash Y$, then the existence of a (local) interpolant for $X \vdash Y$ immediately follows from the inductive hypothesis and the local interpolation property holding for R (cf. Theorem 3.7). $\square$

Applying the previous corollary to D.LE-sequents without structural connectives immediately yields the following

Corollary 3.9 (Lyndon interpolation). *For any LE-language $\mathcal{L}_{LE}$ and any $\mathcal{L}_{LE}$ -sequent $\varphi \vdash \psi$, if $\mathbf{L}_{LE} \models \varphi \vdash \psi$, then $\mathbf{L}_{LE} \models \varphi \vdash \gamma$ and $\mathbf{L}_{LE} \models \gamma \vdash \psi$ for some $\gamma \in \mathcal{L}_{LE}$ such that $\text{Var}^+(\gamma) \subseteq \text{Var}^+(\varphi) \cap \text{Var}^+(\psi)$, and $\text{Var}^-(\gamma) \subseteq \text{Var}^-(\varphi) \cap \text{Var}^-(\psi)$, and γ can be effectively computed.*

4. Generalizing Interpolation to Axiomatic Extensions

The proof strategy of the previous section can be extended to certain axiomatic extensions of $\mathbf{L}_{LE}$ equivalently presented by extending D.LE with appropriate rules that preserve both cut-eliminability and interpolation.

Definition 4.1. *For any LE-language, a special rule in the language of D.LE is an analytic structural rule (see Definition 2.3) of one of the following general forms:*

$$\frac{\Pi \vdash \Psi_1 \quad \dots \quad \Pi \vdash \Psi_n}{\Pi \vdash \Psi_0} \quad \frac{\Phi_1 \vdash \Sigma \quad \dots \quad \Phi_n \vdash \Sigma}{\Phi_0 \vdash \Sigma}$$

where Π (resp. Σ) is an atomic metavariable for structures which does not occur in any Ψ_i (resp. Φ_i) for any $0 \leq i \leq n$.

An interpolation-safe rule is a special rule such that each Ψ_i and Φ_i for any $1 \leq i \leq n$ contain at most one occurrence of at most one metavariable for structures and possibly zeroary structural connectives.

For instance, all analytic-inductive modal reduction principles⁵ can be equivalently captured by interpolation-safe rules. An example is the well-known axiom $\Diamond\Box p \vdash \Box\Diamond p$, which is equivalently captured by the following rule:

$$\frac{X \vdash \Box \blacksquare Y}{X \vdash \blacksquare \Box Y}$$

Another example of this type is the axiom $\Diamond\neg p \vdash \neg\Box p$ of modal fundamental logic (cf. Example 2.1), which is equivalently captured, as discussed in Example 2.5, by the following interpolation-safe rule:

$$\frac{X \vdash \neg \hat{\Diamond} Y}{X \vdash \blacksquare \neg Y}$$

Finally, in a language $\mathcal{L}$, such that for all $f \in \mathcal{F}$ and for all $g \in \mathcal{G}$, $n_f = n_g = 1$ and $\varepsilon_f(1) = \varepsilon_g(1) = 1$, all analytic rules are interpolation safe.

Proposition 4.2. *Every interpolation-safe rule preserves the local Maehara interpolation property, and the interpolant of the conclusion is built from the interpolants in the premises using $\wedge$, $\vee$, or the logical counterpart of structural connectives occurring in the conclusion.*

Proof. We show that if (the instantiations of) the premises of the rule have the Maehara interpolation property, then so does (the instantiation of) the conclusion. Moreover, the interpolant is built from the interpolants of the premises only using the logical counterparts of structural connectives occurring in the conclusion of the rule.

Let us assume a basic LE-logic in an arbitrary but fixed language $\mathcal{L}$. Below, let us consider an instance of the first interpolation-safe structural rule described in Definition 4.1 (the other case is analogous) in the language $\mathcal{L}$:

$$\frac{\Gamma \vdash \Delta_1 \quad \dots \quad \Gamma \vdash \Delta_n}{\Gamma \vdash \Delta_0}$$

Let Y be a substructure of the sequent $\Gamma \vdash \Delta_0$. If Y is a substructure of a metavariable X_0 in $W \vdash T_0$ then Lemma 3.6 gives the desired result. The other possibility is that there exists a substructure $U[X_1, \dots, X_k]$ of $\Pi \vdash \Psi_0$ whose instantiation is Y , where the X_j , $1 \leq j \leq k$, are an exhaustive list of the metavariables U contains. Let's denote by Y^j the substructure of Y that instantiates X_j . Notice, in particular, that Π cannot be one of the X_j . This implies that, since the rule is interpolation safe, each premise $\Pi \vdash \Psi_i$ has at most one instance of one variable among X_j that instantiates a substructure of Y . Let's denote this variable by Z_i , and the substructure of Y it instantiates by Y_i .

By assumption, for the instantiation of each premise, there is an interpolant σ_i , such that $Y^j \vdash^{\varepsilon_{X_j}} \sigma_i$ and $\Gamma \vdash \Delta_i [\sigma_i/Y_i]_k$, where $\varepsilon_{X_j} = 1$ if X_j is an $\mathcal{F}$ -metavariable and $\varepsilon_{X_j} = \partial$ otherwise. Consider $I_j := \{i \mid Z_i = X_j\}$, for $1 \leq j \leq k$. Then it is routine to verify, using the introduction rules for $\wedge$ and $\vee$, that $Y_j \vdash^{\varepsilon_{X_j}} \bigwedge_{i \in I_j} \sigma_i$ and $\Gamma \vdash \Delta_i [\bigwedge_{i \in I_j} \sigma_i/Y_i]_k$. Then applying the rule we obtain $\Gamma \vdash \Delta_0 [U[\bigwedge_{i \in I_1}^{\varepsilon_{X_1}} \sigma_i, \dots, \bigwedge_{i \in I_k}^{\varepsilon_{X_k}} \sigma_i]/Y]$. The formula $u[\bigwedge_{i \in I_1}^{\varepsilon_{X_1}} \sigma_i, \dots, \bigwedge_{i \in I_k}^{\varepsilon_{X_k}} \sigma_i]$ is the interpolant. Since each σ_i satisfies the conditions on the restrictions on variables so do these formulas. Notice in particular that if some $I_j = \emptyset$, then the construction above implies that $\bigwedge_{i \in I_0}^{\varepsilon_{X_j}} \sigma_i$ is $\top^{\varepsilon_{X_j}}$. The fact that added connectives are only those that appear as structural connectives in the conclusion of the rule is immediate. This concludes the proof. $\square$

⁵Given an LE-language $\mathcal{L} = \mathcal{L}(\mathcal{F}, \mathcal{G})$, an *analytic-inductive modal reduction principle* is a sequent $s_1(t_1(p)) \vdash t_2(s_2(p))$ such that s_i and t_i are $\mathcal{L}$ -terms generated as follows:

$$s(\alpha) := \alpha \mid f(s(\alpha)) \quad t(\alpha) := \alpha \mid g(t(\alpha)),$$

with $f \in \mathcal{F}$ and $g \in \mathcal{G}$ such that $n_f = n_g = 1$ and $\varepsilon_f(1) = \varepsilon_g(1) = 1$. These sequents capture a strict subclass of modal reduction principles (cf. [31, Definition 4.18]), namely, those that correspond to analytic rules in display calculi (cf. Proposition 2.4).

Theorem 4.1. *Every analytic axiomatic extension of $\mathbf{L}_{LE}$ whose corresponding rule is a special interpolation-safe rule has the Maehara interpolation property.*

Proof. In view of Theorem 3.7 and Lemma 4.2 what is left to show is that the connectives of the interpolants are in the language of $\mathbf{L}_{LE}$. Notice preliminarily that the conclusion of any analytic rule that corresponds to an axiom of $\mathbf{L}_{LE}$ contains structural connectives whose formula counterparts are part of $\mathbf{L}_{LE}$.

We proceed by induction on the size of the proof. By Lemma 4.2 and the above observation, an application of a non-basic analytic rule is only adding to the interpolant connectives in $\mathbf{L}_{LE}$. Inspecting the cases in Theorem 3.7, reveals that connectives outside $\wedge$ and $\vee$ are added only when an introduction rule is applied, and the connective that is added is the one the rule introduces. This concludes the proof. $\square$

We finish this section by showing that Lyndon interpolation property holds for fundamental logic and its basic modal expansions with unary *tense* connectives [11, 12] (cf. Example 2.1). Indeed, if e.g. the Galois-residuals of $\Box$ and $\Diamond$ in Example 2.1 are present in the language, then $\mathcal{L}^* = \mathcal{L}$. The rules corresponding to all the axioms of these logics are interpolation-safe except for the one corresponding to the LE-axiom $p \wedge \neg p \vdash \perp$, which is, as discussed in Example 2.5

$$R_1 \frac{X \vdash \neg X}{X \vdash Y}$$

Proposition 4.3. *If all connectives in $\mathcal{L}$ are unary except conjunction and disjunction, then the local Maehara interpolation property holds for R_1 . The interpolant is a formula in $\mathcal{L}^*$. Furthermore, if $\Gamma \vdash \Delta[\Sigma]$ is an instance of the conclusion, and the display equivalent sequent $\Sigma \vdash^{\varepsilon\Sigma} \Delta'$ contains only structural connectives from the language of $\mathcal{L}$, then the interpolant of Σ is in the language $\mathcal{L}$.*

Proof. Let $\Gamma \vdash \Delta$ be an instance of the conclusion of R_1 , in which case the assumption is $\Gamma \vdash \neg\Gamma$. If Σ is a substructure of Δ , then clearly $\top^{\varepsilon(\Sigma)}$ is the interpolant, since $\Sigma \vdash^{\varepsilon(\Sigma)} \top^{\varepsilon(\Sigma)}$ is derivable. If $\Gamma = \Pi[\Sigma]$, we can apply R_1 again and obtain $\Pi[\Sigma] \vdash \perp$. Then $\Sigma \vdash^{\varepsilon\Sigma} \Pi'[\perp]$ and so $\Sigma \vdash^{\varepsilon\Sigma} \pi'[\perp]$. Since $\mathbf{L}$ contains only unary connectives $\pi'[\perp]$ contains no variables. Clearly $\pi'[\perp] \vdash^{\varepsilon\Sigma} \Pi'[X]$, and so $\Pi[\pi'[\perp]] \vdash X$, for every structure X . In particular, $\Pi[\pi'[\perp]] \vdash \Delta$. Hence $\pi'[\perp]$ is the interpolant. Indeed, if Π' contains only structural connectives from $\mathbf{L}$ then so does the interpolant. $\square$

The following result is a consequence of the proposition above and of Proposition 4.2 and Theorem 3.7.

Corollary 4.4. *Lyndon interpolation property holds for fundamental logic and its basic modal expansions with tense connectives.*

5. Conclusions and Further Directions

In the present paper, we proved that the Lyndon interpolation property holds for the basic LE-logic over any LE-signature. Our proof is constructive, in the sense that interpolants can be extracted algorithmically from derivations in the corresponding display calculus. Moreover, the method is modular and extendable: we generalized the result to axiomatic extensions of LE-logics captured by augmenting the basic display calculus with suitable analytic structural rules. In particular, we identified a class of *interpolation-safe rules* preserving the local Maehara interpolation property. To illustrate the applicability and modularity of the framework, we established Lyndon interpolation for fundamental logic and its tense modal expansion.

A natural direction for further research concerns the extension of the present framework to *inception display calculi*, recently introduced in Chapter 7 of [32] and further developed in [33]. As discussed earlier on (cf. Proposition 2.4), the expressivity of proper display calculi has been characterized in [13] in terms of a proper subclass of the class of *inductive axioms*, i.e. the largest class of axioms to which the

techniques of algorithmic correspondence theory [34, 7, 35] can be uniformly applied. Inception calculi generalize standard display calculi by allowing derivational nesting through higher-level assumptions, thereby providing analytic proof systems for logics axiomatized by any inductive axioms. Extending the present interpolation method to this setting would require a suitable generalization of the local interpolation property from ordinary analytic structural rules to inception rules. Such an extension would significantly broaden the class of axiomatic extensions for which constructive interpolation methods are available.

More generally, one of the main long-term goals of this line of research is the identification of sufficiently general syntactic conditions on analytic structural rules (and, correspondingly, on the analytic axioms they capture) guaranteeing the validity of various interpolation properties. Ideally, one would aim at a proof-theoretic interpolation meta-theorem for display calculi analogous to Belnap’s celebrated meta-theorem on cut elimination. From this perspective, the class of interpolation-safe rules introduced in the present paper can be regarded as a first step toward a broader structural theory of interpolation for modular proof systems.

Towards this end, taking stock of the proof-theoretic approach developed in Chapter 6 of [32], which is based on the methodology introduced in [36], it would be interesting to reformulate interpolation properties in terms of display-equivalent sequents. We expect that this perspective could lead to a finer-grained proof of Theorem 3.9 in the present paper.

Declaration on Generative AI

During the preparation of this work, the author(s) used LLMs in order to: **Grammar and spelling check, Paraphrase and reword**. After using this tool/service, the author(s) reviewed and edited the content as needed and take(s) full responsibility for the publication’s content.

References

- [1] G. D’Agostino, Interpolation in non-classical logics, *Synthese* 164 (2008) 421–435. URL: <http://www.jstor.org/stable/40271081>.
- [2] W. Fussner, Interpolation in Non-Classical Logics, in: B. ten Cate, J. C. Jung, P. Koopmann, C. Wernhard, F. Wolter (Eds.), *Theory and Applications of Craig Interpolation*, Ubiquity Press, 2026.
- [3] I. van der Giessen, R. Jalali, R. Kuznets, Interpolation in Proof Theory, in: B. ten Cate, J. C. Jung, P. Koopmann, C. Wernhard, F. Wolter (Eds.), *Theory and Applications of Craig Interpolation*, Ubiquity Press, 2026.
- [4] T. Lyon, A. Tiu, R. Goré, R. Clouston, Syntactic interpolation for tense logics and bi-intuitionistic logic via nested sequents, in: M. Fernandez, A. Muscholl (Eds.), *28th EACSL Annual Conference on Computer Science Logic (CSL 2020)*, *Leibniz International Proceedings in Informatics*, 2020, pp. 1–16.
- [5] R. Kuznets, Proving Craig and Lyndon interpolation using labelled sequent calculi, in: L. Michael, A. Kakas (Eds.), *Logics in Artificial Intelligence*, Springer International Publishing, Cham, 2016, pp. 320–335.
- [6] R. Kuznets, Interpolation method for multicomponent sequent calculi, in: *International Symposium on Logical Foundations of Computer Science*, Springer, 2015, pp. 202–218.
- [7] W. Conradie, A. Palmigiano, Algorithmic correspondence and canonicity for non-distributive logics, *Annals of Pure and Applied Logic* 170 (2019) 923–974.
- [8] W. Conradie, A. Palmigiano, C. Robinson, N. M. Wijnberg, Non-distributive logics: from semantics to meaning, in: A. Rezus (Ed.), *Contemporary Logic and Computing*, volume 1 of *Landscapes in Logic*, College Publications, 2020, pp. 38–86.
- [9] G. Greco, P. Jipsen, F. Liang, A. Palmigiano, A. Tzimoulis, Algebraic proof theory for LE-logics, *ACM Trans. Comput. Logic* 25 (2024). doi:10.1145/3632526.

- [10] R. N. Almeida, N. Bezhanishvili, S. Lemal, Superalgamation for modal lattices via non-distributive dualities, 2026. URL: <https://arxiv.org/abs/2602.20380>. arXiv:2602.20380.
- [11] W. Holliday, A fundamental non-classical logic, *Logics* 1 (2023) 36–79.
- [12] W. H. Holliday, Modal logic, fundamentally, in: A. Ciabattoni, D. Gabelaia, I. Sedlár (Eds.), *Advances in Modal Logic*, Vol. 15, College Publications, 2024.
- [13] G. Greco, M. Ma, A. Palmigiano, A. Tzimoulis, Z. Zhao, Unified correspondence as a proof-theoretic tool, *Journal of Logic and Computation* 28 (2018) 1367–1442.
- [14] J. Chen, G. Greco, A. Palmigiano, A. Tzimoulis, Syntactic completeness of proper display calculi, *ACM Transactions on Computational Logic* 23:4 (2022) 1–46.
- [15] D. Van Dalen, *Intuitionistic logic*, The Blackwell guide to philosophical logic (2017) 224–257.
- [16] C. Rauszer, A formalization of the propositional calculus of H-B logic, *Studia Logica* 33 (1974) 23–34.
- [17] M. Gehrke, H. Nagahashi, Y. Venema, A sahlqvist theorem for distributive modal logic, *Annals of pure and applied logic* 131 (2005) 65–102.
- [18] J. M. Dunn, Positive modal logic, *Studia Logica* 55 (1995) 301–317.
- [19] H. P. Sankappanavar, Semi-De Morgan algebras, *The Journal of symbolic logic* 52 (1987) 712–724.
- [20] R. I. Goldblatt, Semantic analysis of orthologic, *Journal of Philosophical logic* (1974) 19–35.
- [21] J. Lambek, On the calculus of syntactic types, *Structure of language and its mathematical aspects* 12 (1961) 166–178.
- [22] N. Galatos, P. Jipsen, T. Kowalski, H. Ono, *Residuated lattices: an algebraic glimpse at substructural logics*, volume 151, Elsevier, 2007.
- [23] M. Moortgat, Symmetric categorial grammar, *Journal of Philosophical Logic* 38 (2009) 681–710.
- [24] J.-Y. Girard, Linear logic, *Theoretical computer science* 50 (1987) 1–101.
- [25] G. Birkhoff, J. Von Neumann, The logic of quantum mechanics, in: *The Logico-Algebraic Approach to Quantum Mechanics: Volume I: Historical Evolution*, Springer, 1975, pp. 1–26.
- [26] M. L. Dalla Chiara, R. Giuntini, et al., Quantum logics, *Handbook of philosophical logic* 6 (2002) 129–228.
- [27] W. Conradie, S. Frittella, A. Palmigiano, M. Piazzai, A. Tzimoulis, N. M. Wijnberg, Categories: how I learned to stop worrying and love two sorts, in: J. Väänänen, Å. Hirvonen, R. de Queiroz (Eds.), *Logic, Language, Information, and Computation*, LNCS 9803, Springer, 2016.
- [28] W. Conradie, A. Craig, A. Palmigiano, N. M. Wijnberg, Modelling informational entropy, in: *Logic, Language, Information, and Computation: 26th International Workshop, WoLLIC 2019, Utrecht, The Netherlands, July 2-5, 2019*, Proceedings 26, Springer, 2019, pp. 140–160.
- [29] S. Maehara, On craig interpolation theorem (in japanese), *Sugaku* 12 (1961) 235–237.
- [30] S. Maehara, G. Takeuti, A formal system of first-order predicate calculus with infinitely long expressions, *Journal of the Mathematical Society of Japan* 13 (1961) 357–370.
- [31] J. van Benthem, Modal correspondence theory, Ph.D. thesis, University of Amsterdam, 1977.
- [32] A. De Domenico, Unified correspondence, proof-theoretically, Phd-thesis - research and graduation internal, Vrije Universiteit Amsterdam, 2025. doi:10.5463/thesis.1412.
- [33] A. De Domenico, G. Greco, A. Palmigiano, Inception display calculi, 2026. To appear in the Special Issue of *Studia Logica* in memory of Nuel Belnap (1930–2024), edited by Thomas Müller, Tomasz Placek, and Heinrich Wansing.
- [34] W. Conradie, A. Palmigiano, Algorithmic correspondence and canonicity for distributive modal logic, *Annals of Pure and Applied Logic* 163 (2012) 338–376.
- [35] W. Conradie, S. Ghilardi, A. Palmigiano, Unified Correspondence, in: A. Baltag, A. Smets (Eds.), *Johan van Benthem on Logic and Information Dynamics*, volume 5 of *Outstanding Contributions to Logic*, Springer International Publishing, 2014, pp. 933–975.
- [36] J. Brotherston, R. Goré, Craig Interpolation in Displayable Logics, in: K. Brunnler, G. Metcalfe (Eds.), *Automated Reasoning with Analytic Tableaux and Related Methods*, Springer Berlin Heidelberg, Berlin, Heidelberg, 2011, pp. 88–103.