

Projective characterization of higher-order quantum transformations

Timothée Hoffreumon and Ognian Oreshkov

Centre for Quantum Information and Communication (QuIC), École polytechnique de Bruxelles, CP 165/59, Université libre de Bruxelles, Avenue F. D. Roosevelt 50, 1050 Brussels, Belgium.

Transformations of transformations, also called higher-order transformations, is a natural concept in information processing, which has recently attracted significant interest in the study of quantum causal relations. In this work, a framework for characterizing higher-order quantum transformations which relies on the use of superoperator projectors is presented. More precisely, working with projectors in the Choi-Jamiołkowski picture is shown to provide a handy way of defining the characterization constraints on any class of higher-order transformations. The algebraic properties of these projectors are furthermore shown to obey rules similar to *multiplicative additive linear logic (MALL)*, providing an intuitive way of comparing any two classes through their projectors. The main novelty of this work is the introduction to the algebra of the ‘prec’ connector. It is used for the characterization of maps that are no signaling from input to output or the other way around. This allows to assess the possible signaling structure of any transformation characterized within the projective framework. The properties of the prec are moreover shown to yield a normal form for projective expressions. This hints towards a general way to compare different classes of higher-order transformations.

1 Introduction

The same way that a quantum channel describes the most general transformation mapping an in-

Timothée Hoffreumon: t.hoffreumon@gmail.com, currently at the Mathematical Institute of the Slovak Academy of Sciences, Bratislava, Slovakia.

put quantum state to an output quantum state [1], a *quantum supermap* describes the most general transformation mapping an input quantum channel to an output quantum channel [2]. Interpreting the quantum channel as a transformation between states, the supermap is then a transformation of transformations. For that reason, it is called *higher-order* transformation. Since nothing forbids *a priori* to nest transformations of transformations, one can consider successive nestings to recursively build whole hierarchies of higher-order transformations [3, 4, 5].

Fragments of a quantum circuit are a concrete instance of the use of a higher-order hierarchy. A fragment of quantum circuit that ‘goes around’ a channel is a supermap: it takes a channel as input and outputs a channel; it is a ‘second-order object’ with respect to seeing the channels as the ‘first-order objects’. This supermap, albeit composed of several interconnected channels, can in turn be seen as a single fragment around which a larger circuit can go. Doing so defines a ‘third-order object’ that can be called a super-supermap: it takes a supermap as input and outputs a channel. And so on. This ensuing hierarchy has been defined under the name *quantum comb formalism* [3], which has proven to be a valuable tool in the field of quantum information theory.

With a different goal than modeling circuit fragments, supermaps with multiple inputs were subsequently studied. First, the *quantum switch* was proposed as a supermap that takes two channels and outputs them in an order that depends on a control qubit [6]. Soon after, *Process Matrices* (PM) were proposed as a general framework of supermaps that take a fixed number of quantum instruments [7] and map them to a joint probability for their outcomes [8]. (In the case when the inputs are channels, that probability is 1.) Both concepts led to the identification of

Indefinite Causal Order (ICO) as a feature of supermaps.

One can then wonder what differentiates the switch from a comb, or the PM from a comb. In particular, why do certain maps and hierarchies of maps feature ICO while others do not? Motivated by these considerations, the goal of this work is to present a framework that formalizes and characterizes higher-order quantum transformations. The two main questions answered are ‘Given an operator on a set of input and output Hilbert spaces, does it represent (the Choi-Jamiołkowski operator of) a higher order object?’ and ‘What is (are) the underlying causal structure(s) of such an object?’. This work extends two previous characterizations, one done using type theory [4, 5], and the other using category theory [9]. This extension relies on the use of superoperator projectors [10, 11, 12, 13].

These projectors have a twofold advantage: first, they make the characterization more straightforward, as one can answer the first question simply by applying the projector corresponding to a given higher-order object on the operator. Second, they offer an intuitive explanation of the type-theoretic semantics of higher order: the algebraic rules for composing these projectors correspond to the semantic rules for forming new types.

The paper and its findings are organized as follows: First, the type-theoretic framework of Ref. [4, 5] is reviewed in Section 2. In Section 3, types are put in correspondence with subsets of operators called *state structures* that are characterized by a superoperator projector. Looking at these projectors instead of the state structures, and, by extension, the types they define, is the idea behind the projective characterization. The type connectives $\{\bar{\cdot}, \otimes, \rightarrow\}$, which respectively correspond to the notions of discarding, parallel composition, and evolution of states, are shown to correspond to rules on how to derive new projectors from the ones characterizing states. It is then proven that this way of characterizing higher-order transformations recovers the type theory. In Section 4, it is shown that the convenient aspect of expressing the type connectives as operations on projectors is that two new operations are obtained ‘for free’: the intersection $\cap$ and union $\cup$ of projectors. This lifts the type system to a particular algebra over $\mathbb{C}$, which is shown in App.

A to be a Boolean lattice $\{\bar{\cdot}, \cap, \cup\}$, equipped with two bilinear compositions $\{\otimes, \rightarrow\}$. Hence, the characterization of higher-order transformations is simplified into manipulations of formulae under Boolean-logic-like rules. In particular, two state structures are equivalent up to normalization if their projectors can be shown equivalent using the rewrite rules of the algebra. It is also observed in this section that the algebra of projectors, when interpreted as a logic, shares similarities with linear logic [14]. In Section 5, the signaling structure of compositions is studied: a third composition of projectors, the ‘prec’ $\prec$, is introduced in the algebra. It encodes the composition of two subsystems in a way that allows signaling in a single fixed direction. It is then shown that the other two, $\otimes$ and $\rightarrow$, correspond to, respectively, no signaling and two-way signaling compositions. In addition to its obvious role of making the underlying signaling structure of a set of transformations apparent, the prec has useful algebraic properties: as presented in Section 6.1, these allow one to rewrite projectors into a normal form, which proves useful for comparing state structures. Finally, in Section 6.2, an example of the use of the projective characterization is given: we recover the result [5, Prop. 6] that quantum combs are quantum networks, which have a single signaling direction.

Three additional examples of the use of the projective characterization are provided in the appendices. First, the characterizations of POVM, channel, and bipartite process matrix formalisms are presented in App. E as an example of what can be done with the results of Sec. 3. Next, a channel theory based on a post-selected state structure is presented in App. F as an example of what can be done in generality with the algebra, according to the results of Sec. 4. Lastly, in App. G, the repeated nesting of channels (map, supermap, and super-supermap) is shown to result in ICO as an illustration of the results presented in Sec. 5 and 6.

2 Types

The formalism of quantum channels and operations [1] describes transformations of quantum states into quantum states, but is unable to effectively describe the transformations whose inputs and outputs are transformations themselves [6].

To overcome this issue, quantum supermaps were introduced [2, 3] and then Perinotti [4] and Bisio [5] extended the concept into an axiomatic framework to classify any transformations of transformations, or *higher-order quantum maps*.

At the core of their work is the utilization of the Choi-Jamiołkowski (CJ) isomorphism [15, 16]. Define $\mathcal{H}^A$ to be a Hilbert space of finite dimension d_A associated to system A , $\mathcal{L}(\mathcal{H}^A)$ to be the space of linear operators on $\mathcal{H}^A$, and $\mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^B))$ to be the space of linear maps from $\mathcal{L}(\mathcal{H}^A)$ to $\mathcal{L}(\mathcal{H}^B)$. Let $|\phi^+\rangle = \sum_{i=0}^{d_A-1} |i\rangle^{A'} \otimes |i\rangle^A$ be a(n unnormalized) maximally entangled state on space $\mathcal{H}^{A'} \otimes \mathcal{H}^A$, with A' a copy of A , and $\phi^+ \equiv |\phi^+\rangle\langle\phi^+|$ its density operator representation. Then,

$$\begin{aligned} \mathcal{M} &\in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^B)), \\ \mathcal{M} &\mapsto M \in \mathcal{L}(\mathcal{H}^{A'} \otimes \mathcal{H}^B): \\ M &\equiv [(\mathcal{I} \otimes \mathcal{M}) \{\phi^+\}]^T, \end{aligned} \quad (1)$$

is an isomorphic bijective mapping. The correspondence sends linear maps between operator spaces to operators in tensor product spaces. It has the properties that a Hermitian-preserving (HP) map is mapped to a Hermitian operator and a completely positive (CP) map is mapped to a positive semi-definite operator. To recover the action of the map, the ‘reverse direction’ of the CJ correspondence is used:

$$\mathcal{M}(V_A) = (\text{Tr}_A [M_{AB} (V_A \otimes \mathbb{1}_B)])^T, \quad (2)$$

where $V_A \in \mathcal{L}(\mathcal{H}^A)$ is an arbitrary operator, and $\mathbb{1}_B$ the identity operator in $\mathcal{L}(\mathcal{H}^B)$. (Capital letters V , N , and U will be reserved to denote operators on single-partite space, while W and M will be for operators on bipartite ones. When needed, subscript letters will be used to clarify to which tensor factor an operator is associated.)

Using this correspondence, maps, maps of maps, and so on, can all be represented as sets of operators on some composite space. These sets resemble formally a constrained version of the set of quantum states on the composite space, and this is how one can compare them: by characterizing the constraints associated with a given set of supermaps.

Then it is possible to define **types** of maps by associating to each of these constrained state

spaces a given *type* [4, 5]. One first defines the base types $A, B, C, \dots$ associated with given systems $A, B, C, \dots$ upon which some parties A(lice), B(ob), C(harlie),... can act¹. A base type, therefore, indicates the state space (in density matrix form) of a given finite-dimensional quantum system as well as its label². (The base types will bear the same label as the system they are associated with. To avoid complicated notation when we describe types instantiated on different systems, we treat them as different, even though two systems A and B could have isomorphic state spaces. If needed, we will indicate by words when two types are isomorphic.)

Out of the base types, one can define transformations as a new type constructed from the base types. For example, if the operator M in Eqs. (1) and (2) above was the CJ representation of a quantum channel, it would transform every element of base type A into an element of base type B , so its type would be ‘of a transformation from type A to type B ’, noted as ‘type $A \rightarrow B$ ’. Other types can be defined by repeating this construction, yielding the hierarchy of higher-order transformations: if types A and B are valid types (not necessarily base types), then $A \rightarrow B$ is a valid type associated with the transformations having an input of type A and an output of type B . The only required notion for defining the hierarchy is that of an *admissible map*, which axiomatizes how the transformations are defined. That is, how (the CJ representations of) the admissible maps between any two types A and B are built from the knowledge of the sets of admissible³ maps forming types A and B .

The formalism is an instance of a *type system*. This “ $\rightarrow$ ” connector, here nicknamed **trans-**

¹Note that Refs. [4, 5] use the terminology ‘elementary type’ instead of ‘base type’, reserve capital letters to base types, and refer to the ‘constrained state spaces’ associated with a given type A as the ‘set of deterministic events of type A ’, which is noted $T_1(A)$.

²Remark that Ref. [9] provides a more general construction by allowing the state space of base types to be different than the one of quantum systems.

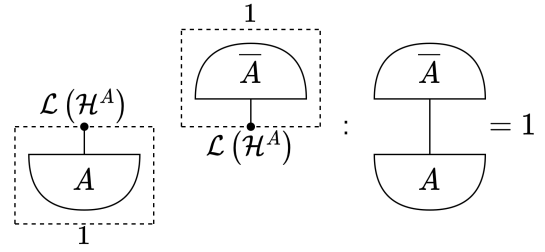
³The admissibility axioms in Ref. [5] essentially generalize the definition of quantum channels by requiring the transformations to be linear maps that preserve the concepts corresponding to complete positivity and trace-preservation for type A and type B . The notion of an admissible map will be formally restated in our formalism under the name of a structure-preserving map, Def. 6.

formation, is the key element of the type theory of higher order transformations: each set can be seen as an abstract type, and new types can be defined out of existing ones using the transformation connector as a semantic rule. Alongside the base types and the transformation, three other symbols must be added to the *alphabet* for writing types: the first one is the trivial type, noted “1”. This is the type associated with a one-dimensional Hilbert space $\mathbb{C}$ whose associated set is made of the number 1, thus corresponding to the absence of a system. The other two symbols are parentheses, “(” and “)”, needed because $\rightarrow$ is not associative. Indeed, $A \rightarrow (B \rightarrow C) \neq (A \rightarrow B) \rightarrow C$ as the first type transforms a ‘state’ of type A into a ‘channel’ of $B \rightarrow C$, whereas the second transforms a ‘channel’ of $A \rightarrow B$ into a ‘state’ of C .

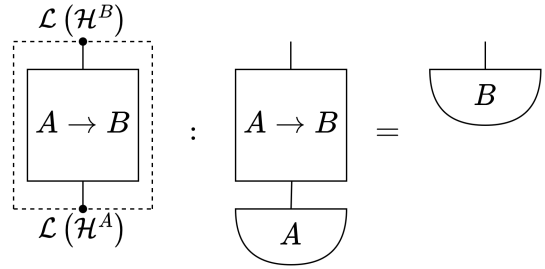
Some graphical notation will be used to represent the types, refer to Fig. 1 in which diagrams are read from bottom to top: open wires are Hilbert spaces of dimension > 1 , and an absence of wire is a space of dimension one (i.e., a trivial system with trivial type 1). A box between two open wires then represents an admissible map (a transformation type). For example, in the left part of subfigure 1b, the box within dashed lines can be seen as (the CJ representation of) any map M of type $A \rightarrow B$; it connects an ‘input’ space (in this case, $\mathcal{L}(\mathcal{H}^A)$) over which the operators of a certain type are defined (in this case, A) to an ‘output’ space over which the operators of another type are defined (in this case, the output space is $\mathcal{L}(\mathcal{H}^B)$ and the type is B). Note that the output type is possibly isomorphic to the input (meaning they are associated with the same set of operators) but they are defined on different spaces nonetheless.

Every type can actually be seen as a transformation from the trivial type to itself, reflecting that any operator can be seen as a linear map from the base field to itself. In the type system, it means that the relation $A = 1 \rightarrow A$ holds. Graphically, it means that a bottom half-circle marked with the letter A is any operator of type A . For example, the left part of subfigure 1a can be seen as any operator $V \in \mathcal{L}(\mathcal{H}^A)$ of type A but also as a transformation of type $1 \rightarrow A$, whence the absence of bottom wire.

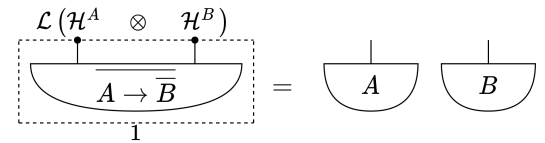
Two special ways of building new types have



(a) An element of type A is represented as a bottom half-circle (left). An element of type $\bar{A} \equiv A \rightarrow 1$ is a top half-circle (center). The closed circuit means that every element of type $\bar{A}$ takes each element of type A to the number 1 (right).



(b) An element of a type $A \rightarrow B$ is represented as a box. The half-closed circuit means that every element of type $A \rightarrow B$ transforms each element of A into one of B .



(c) An element of type $(A \rightarrow (B \rightarrow 1)) \rightarrow 1 = \overline{A \rightarrow B}$ is defined as the parallel composition $A \otimes B$ of two types, so it is represented as its two constituents drawn side by side.

Figure 1: Picturing types graphically. In the following, the notion of a type like A will be subsumed by the one of a state structure $\mathcal{A}$, but the diagrammatic notation will remain the same for simplicity.

their own shorthand notation⁴ as well as graphical depiction because of their interpretation as sets of operators: the ‘functionals on a type A ’, noted $\bar{A} \equiv A \rightarrow 1$, and the ‘tensor product of types A and B ’, noted $A \otimes B \equiv (A \rightarrow (B \rightarrow 1)) \rightarrow 1$. (Why this latter definition corresponds to a notion of tensor product is explained in Ref. [4]; we shall revisit it in this article and explain it using the projector algebra.)

The functional generalizes the notion of discarding to higher-order maps: its type sends a

⁴Such inessential definitions introduced to make the syntax easier to interpret are often called ‘syntactic sugars’ in computer sciences.

type A to the trivial type, meaning it sends the generalization of a state space to the single-element set $\{1\}$. In other words, $\bar{A}$ is the type of the *effects* on A , mapping elements of type A to the number 1.

Such types are represented by a top half-circle, which ‘closes a wire’ and therefore depicts the pairing of a state with an effect to produce a number. This pairing is represented in $\mathcal{L}(\mathcal{H}^A)$ via the inner product so the central part of subfigure 1a can be seen as $\text{Tr}_A [N^\dagger \cdot]$ where $N \in \mathcal{L}(\mathcal{H}^A)$ is of type $A \rightarrow 1 = \bar{A}$. We take as a convention to represent negated types (i.e. with a bar on top of it) as such so they are interpreted as the functional $\text{Tr}_A [N^\dagger \cdot]$ instead of the operator N . This is closer in spirit to how the type reads as ‘taking an element of type A to the number 1’. The right part of the subfigure is then read as the equation $\text{Tr}_A [N^\dagger V] = 1$ for all V of type A and N of type $\bar{A}$. More generally, any closed wire can be seen as applying the reverse direction of the CJ isomorphism and any equation in graphical representation is true for all operators of the types on the left. For example, the right part of subfigure 1b is read as the equation $(\text{Tr}_A [M_{AB} \cdot (V_A \otimes \mathbb{1}_B)])^T = U_B$ for all M and V of types $A \rightarrow B$ and A , respectively, and for some U of type B (the choice of U may depend on the choices of M and V).

Notice that, as operations on types, the ‘bar’ does not commute with the ‘arrow’ i.e., $\overline{A \rightarrow B} \neq \bar{A} \rightarrow \bar{B}$. Indeed, discarding a channel is a different operation from transforming an effect into another one. However, the type $\bar{A} \rightarrow \bar{B}$ of transformations between effects is equivalent to the type $A \leftarrow B$ of reversed transformations between states [5]. (The notation $A \leftarrow B \equiv B \rightarrow A$ will often be used to keep the factors ordered alphabetically. In general, all equivalences of types are up to a reordering of the tensor factors.) This logic-like relation⁵ encodes the fact that the action of a channel from state in $\mathcal{L}(\mathcal{H}^B)$ to a state in $\mathcal{L}(\mathcal{H}^A)$ is equivalent to the action of its adjoint from an effect (represented) in $\mathcal{L}(\mathcal{H}^A)$ to one in $\mathcal{L}(\mathcal{H}^B)$. This relation further implies that the

⁵If the bar is seen as negation and the arrow as implication, this equivalence is indeed the well-known ‘B implies A if and only if not-A implies not-B’.

functionals on functionals are equal to the base type, $\overline{\bar{A}} = A$, i.e. that the ‘bar’ is an involution, since $\bar{A} \rightarrow 1 = A \leftarrow 1$.

The tensor product generalizes the notion of parallel composition to higher-order maps: it takes two types A and B and fuses them into a new bipartite type $A \otimes B$ behaving accordingly. Indeed, $A \rightarrow \bar{B}$ is naively a type that takes an input of a type A and outputs one of type $\bar{B}$. However, this type can also be read in the other direction as $A \rightarrow \bar{B} = \bar{A} \leftarrow B$, so it is equivalently interpreted as a type having an input of type B . Hence, the functionals on such a type, i.e. of type $A \rightarrow \bar{B}$, must be composed of a subsystem of type A and one of type B , since both are valid inputs to $A \rightarrow \bar{B}$. In addition, these subsystems must be somewhat independent of each other as each can be interpreted as the first to be inputted into the objects of type $A \rightarrow \bar{B}$; the choice of an element of A should not influence the choice of one in B and vice-versa. Therefore, type $A \rightarrow \bar{B}$ is a way to combine types A and B in a ‘parallel’ fashion, motivating the notation $A \rightarrow \bar{B} \equiv A \otimes B$. It can then be shown that the set of operators associated with $A \otimes B$ is indeed related to the set spanned by the tensor products of operators of type A with those of type B ⁶. For example, $A \otimes B$ corresponds to the type of all bipartite states in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ if A and B are base types, whereas $(A \rightarrow C) \otimes (B \rightarrow D)$ is the type of all no signaling channels [17, 18] from $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ to $\mathcal{L}(\mathcal{H}^C \otimes \mathcal{H}^D)$; this is the set spanned by all tensor products of channels of type $A \rightarrow C$ with those of type $B \rightarrow D$. Contrastingly, $(A \otimes B) \rightarrow (C \otimes D)$ is the type of all bipartite channels from $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ to $\mathcal{L}(\mathcal{H}^C \otimes \mathcal{H}^D)$ as it transforms bipartite states of type $A \otimes B$ into

⁶For now, this is only an intuitive explanation to motivate the definition of composition from Refs. [4, 5] which was chosen because type $A \otimes B$ for A and B base types precisely corresponds to the state space of a composite quantum system in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$. They then proceed to show that this definition of the tensor product for two arbitrary types results in a suitable notion of the tensor product of their associated state spaces, namely the intersection of the affine span of the tensor products with the cone of positive operators, see Sec. 5.e of Ref. [5]. We will properly restate this definition in Sec. 3.3. In Sec. 5.3, the intuition we presented in the text will be turned into a formal justification for choosing this definition of the tensor product of types; it will be shown to follow only from requirements on the signaling relations.

ones of $C \otimes D$.

Multipartite types like $A \rightarrow B$ are accordingly represented by diagrams with multiple open wires – one for each tensor factor of the underlying Hilbert space. The tensor product of two types is particular as it is represented as the diagrams for its constituents next to each other rather than a joint diagram, see subfigure 1c. Be aware that this graphical representation does not imply that the set of operators of type $A \otimes B$ is separable. For example, a maximally entangled state in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ is of type $A \otimes B$ so it is one of the operators represented by the disjoint diagram on the right of the subfigure.

As expected from a tensor product, this is a commutative and associative composition as $A \otimes B = B \otimes A$ (since $\overline{A \rightarrow B} = \overline{B \rightarrow A}$ up to a reordering of tensor factors) and $A \otimes (B \otimes C) = (A \otimes B) \otimes C$ (this stems from the ‘uncurrying’ property of quantum (super)maps, $A \rightarrow (B \rightarrow C) = (A \otimes B) \rightarrow C$, which encodes the idea that transforming a state into a channel is equivalent to having a channel with two inputs⁷).

Thus, starting from some postulates, a trivial type 1, base types $A, B, C, \dots$, and the “ $\rightarrow$ ” connector rule, all the higher-order generalizations of the quantum formalism can be defined using the type system. These in turn yield the constraints defining the set of operators representing the transformations of a given type [4, 5]. For example, let types A_0 and A_1 be, respectively, the set of input and output quantum states of Alice, who applies some quantum operation in between (subsystems associated with the same party will be differentiated with a numeral index starting at 0, and so will be the tensor factors of the associated Hilbert space as well as the types). Then, one can infer that the set of allowed transformations to which she has access is of type $A_0 \rightarrow A_1$. This simple semantic statement is then translated into constraints to apply on the Hilbert space $\mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$ and yields (the Choi-Jamiołkowski representation of) the set of valid quantum channels for Alice. A more complex example is obtained by recovering the set of bipartite process matrices

⁷However, compared to $A \rightarrow B = \overline{A} \leftarrow \overline{B}$ which is a requirement that comes from the identification of a map with its adjoint, the uncurrying property $A \rightarrow (B \rightarrow C) = (A \otimes B) \rightarrow C$ is rather deduced from the definition of the tensor product; see Sec. 5.e of Ref. [5].

[8], whose type corresponds to a deterministic functional on the local quantum instruments of two parties, say Alice and Bob. Knowing that their local instruments sum up to quantum channels, *i.e.* they belong to types $(A_0 \rightarrow A_1)$ and $(B_0 \rightarrow B_1)$, the set of process matrices is the type that takes a composition of these two types as input and outputs a trivial system. In the semantics, this statement corresponds to forming type $((A_0 \rightarrow A_1) \otimes (B_0 \rightarrow B_1)) \rightarrow 1$, from which the constraints for the characterization directly ensue. According to what has been said so far, the type of bipartite process matrices rewritten as $\overline{(A_0 \rightarrow A_1) \otimes (B_0 \rightarrow B_1)}$ informs us that that they can be seen as the set of functionals (indicated by the bar notation) on the no signaling channels shared between Alice and Bob (of type $(A_0 \rightarrow A_1) \otimes (B_0 \rightarrow B_1)$).

In the following sections, we will develop new characterization methods that improve three aspects of the type system: first, it provides a way to obtain the characterization constraints directly from a type formula. Second, it allows to compare any type of higher-order transformations. This latter feat is key to understanding how ICO arises at higher-orders. Third, it expands the type semantics so to add another connector to represent the type of transformations with a fixed signaling direction. Not only does this also help understanding ICO, but this connector can moreover be used to read the signaling relations that may feature any given kind of higher order directly from its type.

3 Projective characterization of the type theory of higher-order quantum transformations

We start with an example. In quantum theory, the most general representation of a destructive measurement on a state $\rho \in \mathcal{L}(\mathcal{H})$ is done by a POVM [19]: a collection of positive operators $\{E_i\}$, the effects, *resolving* the identity, $\sum_i E_i = \mathbb{1}$. The probability of observing outcome i is given by the *Born rule*:

$$p(i|\rho, \mathbb{1}) = \langle E_i, \rho \rangle = \text{Tr} [E_i^\dagger \rho] . \quad (3)$$

The effects then send each state ρ to a probability $p(i|\rho, \mathbb{1}) \in [0, 1]$ through the Hilbert-Schmidt inner product $\langle \cdot, \cdot \rangle$; they can be seen as *probabilis-*

tic functionals from the state space to a probability, $\langle E_i, \cdot \rangle = \text{Tr} [E_i^\dagger \cdot] : \mathcal{L}(\mathcal{H}) \rightarrow [0, 1]$. These functionals sum up to a *deterministic functional* (or *unit effect*⁸) $\text{Tr} [\mathbb{1} \cdot] : \mathcal{L}(\mathcal{H}) \rightarrow 1$ that sends each state to a probability of 1.

When one considers a non-destructive measurement instead, the most general representation is given by a quantum instrument [7]: a collection of completely-positive (CP) trace-non-increasing maps $\{\mathcal{M}_i\} : \mathcal{L}(\mathcal{H}^A) \rightarrow \mathcal{L}(\mathcal{H}^B)$ *resolving* a quantum channel $\sum_i \mathcal{M}_i = \mathcal{M}$. The probability of observing outcome i together with the (unnormalized) output state $\mathcal{M}_i(\rho) \in \mathcal{L}(\mathcal{H}^B)$ is given by

$$p(i|\rho, \mathcal{M}) = \text{Tr} [\mathcal{M}_i(\rho)] . \quad (4)$$

The similarity with the destructive case is striking in the CJ picture [20, 21]:

$$p(i|W, M) = \langle M_i, W \rangle , \quad (5)$$

where the ‘state’ $W = \rho_A \otimes \mathbb{1}_B \in \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ is destructively measured by a collection of positive operators $\{M_i\}$ resolving a channel M , which is the CJ representation of the instrument.

The two formalisms look similar in the CJ picture, but the ‘effects’ $\mathcal{M}_i$ in the quantum instrument formalism are mappings between the states of the POVM formalism, they are thus higher-order effects. Switching from POVMs to quantum instruments can therefore be seen as a construction of a higher-order theory (this example will be reformulated using the projective characterization in App. G.1).

Notice the common threads that will be generalized in this article: the two formalisms involve ‘states’ (ρ and W , respectively) that are measured by a ‘deterministic functional’ ($\mathbb{1}$ and M), resolved into ‘probabilistic functionals’ (E_i and M_i). What changes between the two situations is that the set of deterministic functionals in the POVM case, $\{\mathbb{1}\}$, has a single element, whereas the one in the quantum instrument case, $\{M\}$, has infinitely many. However, the sets of states and deterministic functionals have a similar structure of subspaces of trace-normalized

positive operators, which we will abstract under the name ‘state structure’. And moreover, the *state structures* of both cases are related: W is the tensor product of a quantum state ρ and the functional on it, $\mathbb{1}$; the set of valid W ’s is thus a state structure obtained by *the tensor product of the state structures* of ρ and $\mathbb{1}$. Likewise, the state structure of the M ’s maps the state structure of ρ in space A to a similar state structure but on space B ; it is *the transformation between two state structures*.

This example motivates the common features of higher-order objects and their projective characterization: the quantum comb formalism [3] or the process matrix formalism [8] are also characterized by a pair of ‘states’ and ‘effects’. Their sets of states are state structures in the CJ picture, as they are sets of positive and normalized operators, each defined over a specific subspace. The projective characterization then consists of finding the projector to that subspace, and classifying how the state structures relate to each other. The three ways of being related that will be formalized in this section appeared in the example: *(deterministic) functional on*, *tensor of*, and *transformation between* state structure(s). Similarly to the types reviewed in the previous section, whose characterization can be inferred from the base type and the rules to form new types, knowing the projector characterizing the ‘base’ state structure will be enough to infer the characterization of every state structure through the rules to form new projectors.

3.1 State structures

Central in this work will be the notion of a *subset of trace-normalized positive elements of an operator system in $\mathcal{L}(\mathcal{H}^A)$* , called a *state structure*. This kind of set is the backbone of the characterization, as to each type of Ref. [4, 5] –or “level in the hierarchy of transformations”– corresponds such a set. The characterization is thus focused on these sets of positive operators that are the images of the types of higher-order maps through the CJ isomorphism, and the abstract structure of these sets is axiomatized through the notion of state structures. A simplification will be to consider only the sets containing an element proportional to the identity operator $\mathbb{1}$; an assumption closely related to the concept of *flatness* defined in Ref. [9]. It is motivated by the fact that the

⁸The nomenclature ‘effect’ and ‘unit effect’ have been used in the literature to refer to the operators respectively associated with the probabilistic and deterministic functionals. In this article, we will treat both nomenclatures interchangeably so that the terms probabilistic and deterministic effects will also appear.

identity is the CJ image of the ‘trace and replace by white noise’ operation; always having an element proportional to the identity then implies that any higher-order map can consist of discarding the input and preparing a purely random output. With states being self-adjoint operators, it is moreover natural to ask that every higher-order map send self-adjoint operators to self-adjoint operators. In the CJ picture, this amounts to considering only self-adjoint Choi matrices. The linear span of a set of self-adjoint operators that contains the identity is called an *operator system*.

Definition 1 (Operator system [22]). *For a given space of operators, an **operator system** is a linear subspace closed under the adjoint that contains the identity.*

We will refer to both an operator system and its subset of trace-normalized positive elements (for a given normalization) by using the calligraphic font of the letter associated with the subsystem it is defined upon (it should be clear from the context which one is under consideration) and we will use primes to differentiate between operator systems defined on the same space. For example, $\mathcal{A}$ and $\mathcal{A}'$ are two different operator systems over $\mathcal{L}(\mathcal{H}^A)$.

Since operator systems are linear subspaces, they can be characterized by linear superoperator projectors noted $\mathcal{P}_A : \mathcal{L}(\mathcal{H}^A) \rightarrow \mathcal{L}(\mathcal{H}^A)$, see Ref. [10, 11, 12]. Projectors on operator systems must be unital and Hermitian-preserving. Because these subspaces are closed, the projectors can be taken self-adjoint without loss of generality.

Definition 2 (Projector on operator system). *A superoperator projector $\mathcal{P}_A : \mathcal{L}(\mathcal{H}^A) \rightarrow \mathcal{L}(\mathcal{H}^A)$ obeying the following conditions for all $V, N \in \mathcal{L}(\mathcal{H}^A)$:*

$$\mathcal{P}_A \{ \mathbb{1} \} = \mathbb{1} ; \quad (6a)$$

$$\mathcal{P}_A \{ V^\dagger \} = (\mathcal{P}_A \{ V \})^\dagger ; \quad (6b)$$

$$\text{Tr} [N^\dagger \mathcal{P}_A \{ V \}] = \text{Tr} [(\mathcal{P}_A \{ N \})^\dagger V] ; \quad (6c)$$

*is called a **projector on (an) operator system**.*

In addition, higher-order transformations should preserve the positivity of the spectrum.

This has to do with the probabilistic interpretation of state structures that have been obtained from a set of linear maps that underwent the CJ isomorphism (this will be explained in Sec. 3.4 below). Briefly put, the state of a party cannot have negative eigenvalues after a transformation, as this would correspond to negative outcome probabilities. In addition, following the ‘admissibility criterion’ of Ref. [5], any higher-order transformation should, in general, be definable on a subsystem. Preservation of the positivity of the spectrum even when acting on a subsystem requires transformations to preserve complete positivity. In the CJ picture, this translates into the operators being positive semi-definite. In addition, as the probabilities must still sum up to one after any higher-order transformation and since positivity is preserved, the transformations can at most rescale the traces by a positive factor. In the CJ picture, this translates into the operators having a fixed trace norm.

These two requirements restrict the operator system to trace-normalized positive operators, meaning that any of its elements V_A have to respect the following set of conditions:

$$V_A \geq 0 , \quad (7a)$$

$$\text{Tr} [V_A] = c_A , \quad (7b)$$

$$\mathcal{P}_A \{ V_A \} = V_A , \quad (7c)$$

where $\mathcal{P}_A$ satisfies conditions (6). The abstract structure they define will play a central role in the characterization of higher-order quantum transformations. We name it a *state structure*.

Definition 3 (State structure). *A **state structure** $\mathcal{A} \in \mathcal{L}(\mathcal{H}^A)$ is a set obeying constraints (7).*

As mentioned in the introductory example, an instance of a state structure is the set of quantum states in density matrix form, $\rho \in \mathcal{L}(\mathcal{H})$, characterized by

$$\rho \geq 0 , \quad (8a)$$

$$\text{Tr} [\rho] = 1 , \quad (8b)$$

$$\mathcal{I}\{\rho\} = \rho . \quad (8c)$$

Here, $\mathcal{I}$ is the *identity mapping*, defined as

$$\forall V \in \mathcal{L}(\mathcal{H}^A) : \mathcal{I}_A(V) = V . \quad (9)$$

The elements of a POVM are actually operators summing up to (or *resolving*) the single-element state structure defined by

$$\mathbb{1} \geq 0, \quad (10a)$$

$$\text{Tr} [\mathbb{1}] = d, \quad (10b)$$

$$\mathcal{D}\{\mathbb{1}\} = \mathbb{1}, \quad (10c)$$

where $\mathcal{D}$ is the *depolarizing superoperator*, obeying

$$\forall V \in \mathcal{L}(\mathcal{H}^A) : \mathcal{D}_A(V) = \frac{\mathbb{1}^A}{d_A} \otimes \text{Tr}_A[V], \quad (11)$$

which projects onto the span of the identity.

Remark that the projectors $\mathcal{I}$ and $\mathcal{D}$ commute. On a given Hilbert space, we will assume all projectors under consideration to always commute pairwise. The reasons behind this assumption will become clearer in the following. Assuming that all projectors under consideration commute pairwise allows to use the freedom in the choice of basis for the CJ isomorphism to ensure that all of them commute with the transposition appearing in Eqs. (1) and (2) when acting on self-adjoint operators. As a consequence, the operator systems under consideration will also be closed under this transposition. (If needed, the reader can refer to App. A.1 for an extended discussion of what is meant and assumed when referring to ‘projectors on operator systems’ in this article. See in particular Eqs. (84) for an abstract formulation of Eqs. (6), and Eqs. (94) for the properties assumed about every set of projectors on operator systems.)

In order to compute probabilities on a state structure, it becomes interesting to consider the *probabilistic families resolving* its elements. This generalizes the idea of POVM elements in quantum theory, for which each POVM is a different resolution of the identity $\sum_i E_i = \mathbb{1}$, as well as the idea of quantum instrument formalism, for which each quantum instrument is a particular family of CP trace-non-increasing maps $\mathcal{M}_i$ providing a resolution of a particular CPTP map $\sum_i \mathcal{M}_i = \mathcal{M}$. Formally,

Definition 4 (Resolution of a state structure). *Let $\mathcal{A}$ be a state structure in $\mathcal{L}(\mathcal{H}^A)$. The set of operators **resolving** $\mathcal{A}$ is the set of all collections of positive operators summing up to an element of $\mathcal{A}$. That is, a set of operators $\{E_i\}$*

*is a **resolution** of $\mathcal{A}$ if and only if*

$$\begin{aligned} \forall E_i \in \{E_i\} \subset \mathcal{L}(\mathcal{H}^A), \exists V \in \mathcal{A} : \\ E_i \geq 0, \\ E_i + \sum_{j \neq i} E_j = V. \end{aligned} \quad (12)$$

3.2 Functionals

Type A has been generalized into state structure $\mathcal{A}$, we now want to generalize the notion of a set of *deterministic effects* from the type $\bar{A}$ to a state structure. It is the set of operators N representing all the functionals mapping each element V of a state structure to the number 1 via the Hilbert-Schmidt inner product for self-adjoint operators, $\text{Tr}[N \cdot V] = 1$ (note that this inner product features a dagger on the left member, which we omit since all operators are positive). In the following, we will always mean deterministic effect when we use the word ‘effect’ alone.

In a previous work [11], it was noticed that since both the set of states and effects must contain an element proportional to the identity, the two subspaces on which the states and effects are respectively defined must be *quasi-orthogonal*⁹. This means that operators V and N must be orthogonal everywhere but at the span of the identity. If $\bar{\mathcal{P}}_A$ is the projector for $\bar{\mathcal{A}}$, it is related to $\mathcal{P}_A$ by $\bar{\mathcal{P}}_A \equiv \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A$. This leads to the rephrasing in terms of projectors of the characterization of $\bar{\mathcal{A}}$ from type theory [5, Lemma 4].

Proposition 1 (Functional). *Let $\mathcal{A}$ be a state structure. Let $\{E_i\}$ a resolution of an element of $\mathcal{A}$ as in Def. 4. Let $\bar{\mathcal{A}}$ be the set of operators taking each element V of $\mathcal{A}$ to the number 1 through the inner product,*

$$N \in \bar{\mathcal{A}} \Rightarrow \forall V \in \mathcal{A} : \text{Tr}[N \cdot V] = 1, \quad (13)$$

and taking each element E_i of every resolution of V to a positive number between 0 and 1, i.e.,

$$\begin{aligned} \forall V \in \mathcal{A}, \forall \{E_i\} : E_i \geq 0, \sum_i E_i = V, \\ \text{Tr}[N \cdot E_i] \in [0, 1]. \end{aligned} \quad (14)$$

⁹The terminology “quasi-orthogonal” originated in the study of maximally Abelian subalgebras (MASAs), see Ref. [23]

Then $\overline{\mathcal{A}}$ is a state structure characterized by the following conditions:

$$N \in \overline{\mathcal{A}} \iff N \geq 0, \quad (15a)$$

$$\text{Tr}[N] = \frac{d_A}{c_A} \equiv c_{\overline{A}}, \quad (15b)$$

$$\{\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A\}(N) \equiv \overline{\mathcal{P}}_A\{N\} = N. \quad (15c)$$

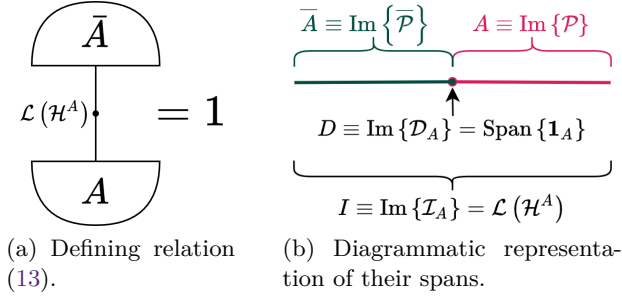


Figure 2: Picturing the characterizing equations of an effect state structure, Proposition 1. Fig. 2a is the graphical representation of the defining relation (13). Fig. 2b is a diagrammatic representation of how the Hilbert space splits between the operator system $\mathcal{A}$ (pink) and $\overline{\mathcal{A}}$ (green). (In the diagram, the regular fonts A and $\overline{A}$ have been used instead of script-style because of software limitations; ‘Im’ means the image of a linear map). Together they span the full space (whole line), yet they only intersect at identity (central dot).

The proof is provided in App. D.1, see Fig. 2a for an illustration. This proposition was first obtained as a theorem in our earlier work on Multi-round Process Matrices (MPM) [11]. A concrete example of a state structure characterized by Prop. 1 are Eqs. (10) obtained from Eqs. (8). Other examples feature the quantum testers [3] whose characterization follows from the one of quantum combs using Prop. 1, and process matrices [8], whose characterization follows from the one of the tensor product of channels. The projective characterizations of single and bipartite process matrices are respectively presented in App. E.3 and E.5.

Notice how Prop. 1 splits the space of operators: the projectors $\mathcal{P}_A$ and $\overline{\mathcal{P}}_A$ obey the relation

$$\mathcal{I}_A = \mathcal{P}_A + \overline{\mathcal{P}}_A - \mathcal{D}_A; \quad (16)$$

this means that the space of all operators is divided between operators V obeying $\mathcal{P}_A\{V\} = V$ and operators N obeying $\overline{\mathcal{P}}_A\{N\} = N$; the only operators X allowed to live on the intersection of

the two subspaces are those obeying $\mathcal{D}_A(X) = X$ since $\mathcal{P}_A \circ \overline{\mathcal{P}}_A = \mathcal{D}_A$. Consequently, the only elements in $\mathcal{A} \cap \overline{\mathcal{A}}$ are multiples of the identity. See Fig. 2b for a diagrammatic illustration of these relations.

Here, the bar in “ $\overline{\mathcal{P}}_A \equiv \mathcal{I} - \mathcal{P}_A + \mathcal{D}$ ” can be seen as an operation at the level of projectors that was applied on $\mathcal{P}_A$. As such, the projector $\overline{\mathcal{P}}_A$ will be called the **negation** of $\mathcal{P}_A$ because the operation has properties similar to a logical *not*; see App. A.3 for the details. This algebraic interpretation of projectors will be developed in Sec. 4 below.

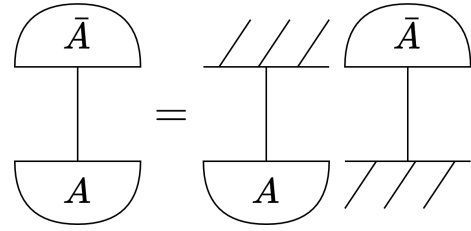


Figure 3: Quasi-orthogonality condition (17) is equivalent to the factorization (18) of the inner product. Graphically, the element of a state structure proportional to $\mathbb{1}$ is by convention represented by the symbol $///$ (this is to single it out from arbitrary elements, which are represented by a half-circle). This diagram should then be read as $\forall V \in \mathcal{A}, \forall N \in \overline{\mathcal{A}}, \text{Tr}[N \cdot V] = \text{Tr}\left[\frac{1}{c_A} \cdot V\right] \text{Tr}\left[N \cdot \frac{1}{c_{\overline{A}}}\right]$. The diagrammatic translation of quasi-orthogonality between two state structures is thus the possibility of ‘cutting’ the wire between them; this disconnection conveys the idea that these two state structures cannot influence each other.

Quasi-orthogonality implies the following property [23, Theorem 2.37 iii)]:

$$\forall V \in \mathcal{A}, \forall N \in \overline{\mathcal{A}}, \text{Tr}[N \cdot V] = \frac{1}{d_A} \text{Tr}[N] \text{Tr}[V]. \quad (17)$$

Combining the trace conditions (7b) and (15b) with equation (17) yields the following.

$$\text{Tr}[N \cdot V] = \text{Tr}\left[N \cdot \frac{1}{c_{\overline{A}}}\right] \text{Tr}\left[\frac{1}{c_A} \cdot V\right]. \quad (18)$$

Fig. 3 provides a graphical interpretation of this relation: a wire connecting a state V with a ‘deterministic functional’ (a trace with unit effect N ; $\text{Tr}[N \cdot V]$) can be ‘cut’ in two pieces (two traces, $\text{Tr}[N \cdot \bullet] \text{Tr}[\bullet \cdot V]$) by adding two ‘discarding effects’ on each loose end ($\frac{1}{c_{\overline{A}}}$ and $\frac{1}{c_A}$ replace their respective $\bullet$). This hints the physical content of Proposition 1: the generalized Born rule is independent of *which particular* deterministic state

V and effect N are getting connected together; both the state and the effect can be replaced by the identity – or any other element of their respective state structures – without altering the outcome probability of 1. Put differently, even when replacing either the state or the effect by pure noise (that is, the discarding effect), it is always a certainty that *something has happened* at the end of the protocol. The connection between quasi-orthogonality and the independence between the choice of (deterministic) state and effect is what allows to generalize the notion of *no signaling* to state structures. This will be developed in Sec. 5.

3.3 Tensor product

The tensor product of two types as a parallel composition is an abstraction that encompasses both the notions of *bipartite quantum state* as well as *no signaling channel* [17, 18]. In terms of projectors, the characterization relies on raising the tensor product operation at the level of superoperators in a natural fashion. Let $\mathcal{P}_A$ and $\mathcal{P}_B$ be respectively projectors on operator systems $\mathcal{A}$ and $\mathcal{B}$. Then $\mathcal{P}_A \otimes \mathcal{P}_B$ acts on $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ so that $\forall q_i \in \mathbb{C}, V_i \in \mathcal{L}(\mathcal{H}^A), U_i \in \mathcal{L}(\mathcal{H}^B)$,

$$(\mathcal{P}_A \otimes \mathcal{P}_B) \left\{ \sum_i q_i (V_i \otimes U_i) \right\} \equiv \sum_i q_i (\mathcal{P}_A \{V_i\} \otimes \mathcal{P}_B \{U_i\}) . \quad (19)$$

This straightforward definition of the tensor product of two projectors on operator systems results in a projector on operator system in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ (see App. A.4.1). Hence $\mathcal{P}_A \otimes \mathcal{P}_B$ is an operation on projectors $\mathcal{P}_A$ and $\mathcal{P}_B$ that will be called their **tensor product** (or tensor in short). A diagrammatic depiction of its span with respect to the bipartition of the Hilbert space is given in Fig. 4.

Definition 5 (No signaling (bipartite) composition). *Let $\mathcal{A}$ and $\mathcal{B}$ be two state structures as in Eqs. (7). Their no signaling composition $\mathcal{A} \otimes \mathcal{B} \subset \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ is the set of all operators*

W characterized by the following constraints:

$$W \in \mathcal{A} \otimes \mathcal{B} \iff W \geq 0, \quad (20a)$$

$$\text{Tr}[W] = c_A c_B, \quad (20b)$$

$$(\mathcal{P}_A \otimes \mathcal{P}_B) \{W\} = W. \quad (20c)$$

Consequently, $\mathcal{A} \otimes \mathcal{B}$ is the intersection of the affine span of tensor products of operators from $\mathcal{A}$ and $\mathcal{B}$ with the cone of positive operators.

Def. 5 relies on the observation that the positive elements of the affine hull of the tensor product of state structures form a state structure [5, Lemma 5] to define a way to compose two state structures. It recovers reasonable notions of parallel composition: the set of bipartite states is obtained by setting $\mathcal{A}$ and $\mathcal{B}$ to be the state structures of quantum states, whereas the set of no signaling (sometimes called *causal* [17]) bipartite channels is obtained by setting $\mathcal{A}$ and $\mathcal{B}$ to be the state structures of channels. (These two examples are respectively presented in App. E.2 and E.4.) We have called it ‘the no signaling composition’ for reasons that will become clear in Sec. 5. For now, this definition will help us for the characterization of the CJ representation of transformations between state structures.

Notice that the decomposition (16) of the identity projector in terms of $\mathcal{P}_A$ and $\bar{\mathcal{P}}_A$ can be generalized to the identity projectors on two systems. Since $\mathcal{I}_{AB} = \mathcal{I}_A \otimes \mathcal{I}_B$ by linearity, taking the tensor product of Eq. (16) with itself gives

$$\begin{aligned} \mathcal{I}_A \otimes \mathcal{I}_B = & \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{P}_A \otimes \bar{\mathcal{P}}_B + \bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B + \bar{\mathcal{P}}_A \otimes \mathcal{P}_B \\ & - \mathcal{P}_A \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \bar{\mathcal{P}}_B - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \mathcal{P}_B \\ & + \mathcal{D}_A \otimes \mathcal{D}_B. \end{aligned} \quad (21)$$

This relation is diagrammatically represented in Fig. 4; it implies that the space is divided between four main parts depending on $\mathcal{P}_A$, $\mathcal{P}_B$, and their negations. Each of these parts might share a common ‘border’, which can be non-trivial, e.g., $\mathcal{P}_A \otimes \mathcal{D}_B$ is at the intersection of $\mathcal{P}_A \otimes \mathcal{P}_B$ and $\mathcal{P}_A \otimes \bar{\mathcal{P}}_B$ since $(\mathcal{P}_A \otimes \mathcal{P}_B) \circ (\mathcal{P}_A \otimes \bar{\mathcal{P}}_B) = \mathcal{P}_A \otimes \mathcal{D}_B$, or not, e.g., $\mathcal{D}_A \otimes \mathcal{D}_B$ is at the intersection of $\mathcal{P}_A \otimes \mathcal{P}_B$ and $\bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B$ since $(\mathcal{P}_A \otimes \mathcal{P}_B) \circ (\bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B) = \mathcal{D}_A \otimes \mathcal{D}_B$. In this latter case, the ‘border’ is trivial as the only elements contained in it are proportional to the identity operator $\mathbb{1}_A \otimes \mathbb{1}_B$. The systematic way to discuss the splitting of the

space of operators induced by the projectors and how to find the ‘borders’, i.e. to obtain intersections of operator systems as compositions of projectors, will be presented in Sec. 4.

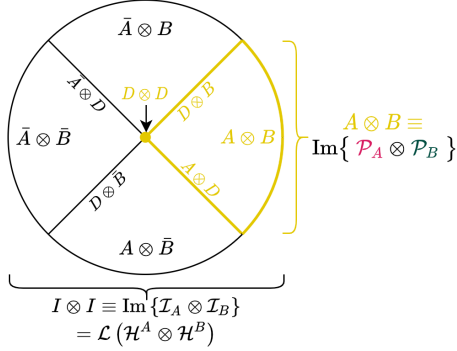


Figure 4: Diagrammatic representation of the image of the tensor product of projectors, as in Def. 5. The full wheel represents $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ while its parts represent its tensor factors, Eq. (21). (Like in Fig. 2b, ‘Im’ means ‘Image’ and e.g. D is a shortcut for $\text{Im}\{D\} = \text{Span}\{1\}$). By convention, the projectors acting on $\mathcal{L}(\mathcal{H}^A)$ are always on the left-hand side of tensor products and vice-versa for B , so that subscripts can be omitted. Note that the intersections are well defined; for example the line ‘ $A \otimes D$ ’, which represents $\text{Im}\{P_A \otimes D_B\}$ is indeed the intersection $A \otimes B \cap A \otimes \bar{B}$.

3.4 Transformations

In the type theory of higher-order transformations, every type is an instance of a transformation, so the previous characterization of functionals (type $\bar{A}$) and bipartite states (type $A \otimes B$) are but special cases of a more general rule. This rule says that if A and B are types, then $A \rightarrow B$ is itself a type that takes elements of A to B . A state of type A is actually a *transformation* of the trivial type – the number 1 – into A , noted $1 \rightarrow A$. Accordingly, the effect of type $\bar{A}$ is the transformation of type A into the trivial type, $A \rightarrow 1$. No signaling composition is the type $A \otimes B \equiv A \rightarrow \bar{B}$.

As we have seen, the set of elements of a given type corresponds to a state structure, ultimately determined by its projector. Thus, to type $A \rightarrow B$ corresponds a state structure, thereafter noted $\mathcal{A} \rightarrow \mathcal{B}$, which is obtained by combining state structures $\mathcal{A}$ and $\mathcal{B}$ in a certain way. This structure should then be characterized by a composite projector built by combining P_A and P_B using the two previously introduced rules.

Following previous approaches [1, 4, 5, 24], it should be possible to randomize the dynamics – to pick a transformation at random – without losing the probabilistic interpretation, implying that the transformation must conserve the total probability as well as probabilistic mixtures. And it should be possible to do so even when acting locally on a tensor factor, i.e. even when it acts on a single part of a no signaling composite state structure [5, Definition 7]. The preservation of probabilistic mixtures implies that the maps representing such dynamics are linear (see Ref. [1, 8, 4, 5] for the proof). It should also be possible to consider entangled input and output pairs [2, 3, 4, 5] –i.e., to have correlated inputs and outputs obtained via the exchange of a side system – without losing the probabilistic interpretation. This generalizes the idea of a CP map and implies that the transformation must be normalized on any state in the tensor product of its input and output and any resolution thereof. This last bit is important as it constrains the CJ representations of transformations to positive semi-definite operators (see Ref. [2, 5]).

Definition 6 (Structure-preserving map). *Let $\mathcal{M}$ be a map from $\mathcal{L}(\mathcal{H}^A)$ to $\mathcal{L}(\mathcal{H}^B)$. This map is **structure-preserving** between $\mathcal{A}$ and $\mathcal{B}$ if*

- 1) *it is linear;*
- 2) *it maps any element of state structure $\mathcal{A}$ to one in $\mathcal{B}$, that is,*

$$\forall V \in \mathcal{A}, \quad \mathcal{M}(V) \in \mathcal{B}; \quad (22)$$

3) *its CJ operator is positive.*

Proposition 2 (Transformation between state structures). *Let $\mathcal{M} \in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^B))$ be a structure-preserving map between state structures $\mathcal{A}$ and $\mathcal{B}$. Call $M \in \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ the Choi-Jamiołkowski representation of this map, as in Eq. (1). Then, the set $\{M\}$ of all such maps belongs to the state structure $\mathcal{A} \rightarrow \mathcal{B}$ defined by the following conditions:*

$$M \in \mathcal{A} \rightarrow \mathcal{B} \iff M \geq 0, \quad (23a)$$

$$\text{Tr}[M] = c_A^{-1} c_B = \frac{c_B}{c_A} d_A, \quad (23b)$$

$$(P_A \rightarrow P_B) \{M\} \equiv (\overline{P_A \otimes P_B}) \{M\} = M. \quad (23c)$$

The proof is presented in App. D.2; see Fig. 5 for an illustration. An example of a state structure characterized by this proposition is the set

of quantum channels, obtained by setting $\mathcal{A}$ and $\mathcal{B}$ to be state structures of quantum states (the examples of single and bipartite channels are respectively presented in App. E.3 and E.4; note that App. E.3 is but the introductory example of this section rephrased in the formalism developed in this paper). Another example is the set of quantum n -combs [3] in which $\mathcal{A}$ is the set of $(n-1)$ -combs and $\mathcal{B}$ the one of channels (this example is revisited in Sec. 6.2).

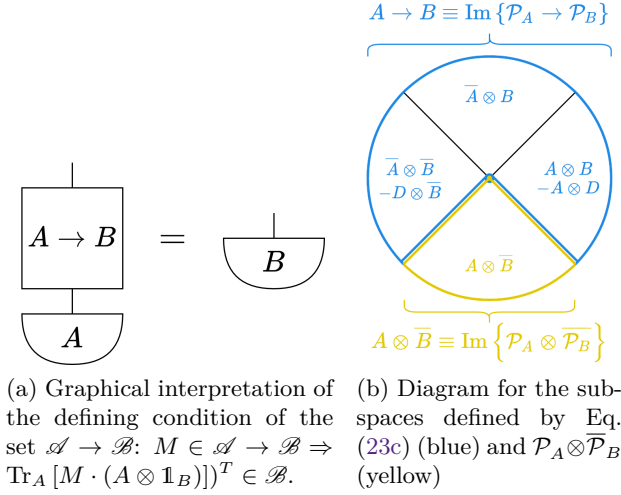


Figure 5: Picturing the characterizing equations of a transformation between state structures, Proposition 2. Fig. 5a is the graphical representation of the defining relation (22). Fig. 5b is a diagrammatic representation of how the Hilbert space splits between the span of the state structure of transformations, $\mathcal{A} \rightarrow \mathcal{B}$ (blue), and the span of the deterministic functionals on it, $\mathcal{A} \otimes \overline{\mathcal{B}}$ (yellow). As in Fig. 2b, they span the full space together (i.e., they cover the whole disk), yet they only intersect at the identity (central dot). Notice how the transformation does not contain the boundaries $\text{Im}\{\mathcal{D}_A \otimes \overline{\mathcal{P}}_B\}$ and $\text{Im}\{\mathcal{P}_A \otimes \mathcal{D}_B\}$ as they belong to $\mathcal{A} \otimes \overline{\mathcal{B}}$ already.

In Proposition 2, a new operation on the projectors has been defined, the **transformation** $\rightarrow$. It combines the tensor and negation operations, $\mathcal{P}_A \rightarrow \mathcal{P}_B \equiv \overline{\mathcal{P}_A \otimes \overline{\mathcal{P}}_B}$ and its full form is the following projector:

$$\mathcal{P}_A \rightarrow \mathcal{P}_B \equiv \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{I}_B + \mathcal{P}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (24)$$

This is a connector whose direction matters. As we did with types, we define a reversed symbol, $\leftarrow$, by $\mathcal{P}_A \leftarrow \mathcal{P}_B \equiv \overline{\mathcal{P}_A \otimes \overline{\mathcal{P}}_B}$ in order to use it without changing the ordering of the tensor factors.

Remark. The above derivation is very close in spirit to Ref. [24] where the projective characterization of transformations was first made in the context of transformation between process matrices, but the obtained projector was missing two terms¹⁰. After verification, it appears that this omission has not hindered the other conclusions of this article [25].

3.5 The algebra of projectors recovers the type theory

We started from the type theory of higher-order transformations for which $\overline{A}$, $A \otimes B$, and $A \rightarrow B$ are the new types that one can obtain from base types A and B under the semantics $\{1, (,), \rightarrow\}$.

We observed that types define *state structures*, which are subsets of operator systems made of trace-normalized positive elements. These state structures are the sets of CJ operators representing higher-order maps for which $\overline{\mathcal{A}}$, $\mathcal{A} \otimes \mathcal{B}$, and $\mathcal{A} \rightarrow \mathcal{B}$ are the new state structures that one can obtain from bases $\mathcal{A}$ and $\mathcal{B}$.

We thus refined the type theory by relaxing the requirement that the base types are quantum states into the base type being a valid state structure. As these state structures are defined on linear subspaces, most of their properties are encoded by the superoperator projector that defines them. These projectors are an abstraction of states structures for which $\overline{\mathcal{P}}_A$, $\mathcal{P}_A \otimes \mathcal{P}_B$, and $\mathcal{P}_A \rightarrow \mathcal{P}_B$ are the new projectors that one can obtain from bases $\mathcal{P}_A$ and $\mathcal{P}_B$. In that respect, the projective characterization amounts to the three algebraic rules $\{\overline{\cdot}, \otimes, \rightarrow\}$ defined in Eqs. (15c), (20c), and (23c).

Lemma 1. *Let $\mathcal{P}_A$ be an arbitrary projector on an operator system (see Def. 2) and acting on space $\mathcal{L}(\mathcal{H}^A)$. Let $\mathcal{P}_B, \dots, \mathcal{P}_K$ be similarly defined projectors, respectively acting on $\mathcal{L}(\mathcal{H}^B), \dots, \mathcal{L}(\mathcal{H}^K)$. Then, every superopera-*

¹⁰What the authors had obtained in the context of process matrix to process matrix transformation was $\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{I}_B + \mathcal{P}_A \otimes \mathcal{P}_B$, where $\mathcal{P}_A$ is the projector characterizing the input process matrix state structure and $\mathcal{P}_B$ the output. Compared to equation (23c), here written fully without negation, $\overline{\mathcal{P}_A \otimes \overline{\mathcal{P}}_B} = \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{I}_B + \mathcal{P}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B$, one sees that two terms are missing. In particular, they are the ones that forbid to postselect on a process matrix before preparing a new one that is purely noisy.

tor acting on $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B \otimes \dots \mathcal{H}^K)$ obtained by combining these projectors using any permutation of the rules $\{\neg, \otimes, \rightarrow\}$ is a projector on operator system.

Proof. Each rule individually results in a projector on an operator system as can be shown by direct computation, meaning that if $\mathcal{P}_A$ and $\mathcal{P}_B$ are projectors on operator systems, then so are $\overline{\mathcal{P}}_A$, $\mathcal{P}_A \otimes \mathcal{P}_B$, and $\mathcal{P}_A \rightarrow \mathcal{P}_B$, see Appendices A.3 and A.4.1 for the full proof of this statement. It remains to show that any combination of these rules are also giving projectors also obeying Eqs. (6). Using that $\mathcal{P}_A \rightarrow \mathcal{P}_B = \overline{\mathcal{P}}_A \otimes \overline{\mathcal{P}}_B$, any expression is either an overall negation of a projector, $\overline{\mathcal{P}}$, or it is an overall tensor of k factors, $\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K$. In the first case, it is a valid projector provided that the projector being negated is valid. In the second case, the k factors can be seen as a tensor product of two factors using the associativity of the tensor, $(\mathcal{P}_A \otimes \dots) \otimes (\mathcal{P}_K)$, which is a valid projector provided $\mathcal{P}_A \otimes \dots$ and $\mathcal{P}_K$ are. This reasoning repeats inductively until one reaches each individual projector $\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K$ which are projectors on operator systems by assumption. $\square$

And because of that, they define valid state structures globally.

Corollary 1. *Any set build from state structures using the rules $\{\neg, \otimes, \rightarrow\}$ (characterized by, respectively, Prop. 1, Def. 5, and Prop. 2) is a state structure itself.*

Thus, when a set $\{\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K\}$ of k ‘base’ projectors has been defined and associated with k systems $A, B, \dots, K$, one can form arbitrary combinations (for example $(\mathcal{P}_A \rightarrow \mathcal{P}_B) \otimes \overline{\mathcal{P}}_F$) and this is guaranteed to define a new composite state structure on the joint space once the normalizations have been fixed (in the example $(\mathcal{A} \rightarrow \mathcal{B}) \otimes \overline{\mathcal{F}} \subset \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B \otimes \mathcal{H}^F)$ is a valid state structure once the constants c_A , c_B , and $c_{\overline{F}}$ are fixed). Now, the characterization has shown that the obtained rules are redundant: instead of using $\{\neg, \otimes, \rightarrow\}$, one can use $\{\neg, \rightarrow\}$ only, as Eq. (23c) shows that any occurrence of the no-signaling composition can be replaced by a combination of the negation and the transformation. Similarly, instead of using $\{\neg, \rightarrow\}$, one can use the transformation $\rightarrow$ only. To do so one can

leverage the trivial system 1, which is the single-element state structure composed of the number one and characterized by a projector which is also the number one (which is the only projector on $\mathbb{C}$). Using such a rule, one can see that $\overline{\mathcal{P}}_A = \mathcal{P}_A \rightarrow 1$ since $\mathcal{P}_A \rightarrow 1 = \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A$ using Eq. (24). Hence, the operations on projectors can be reduced to only the transformation, and together with the parentheses and the trivial projector 1, they give the same semantics $\{1, (,), \rightarrow\}$ as the type theory of Refs. [4, 5]. This is not a coincidence: when the base state structures are chosen to be the state spaces of quantum states, i.e., when the base projectors are chosen to be the identity projector $\mathcal{I}$ and the associated traces are set to 1, the type formalism is recovered. The next theorem is a technically precise version of this statement, given for completeness only and aimed at a reader familiar with Ref. [5].

Proposition 3. *When the base state structures are set to be the sets of density matrices, the state structures characterized by the projective characterization under the rules $\{\neg, \otimes, \rightarrow\}$ are in one-to-one correspondence with the sets of ‘deterministic events’ characterized using the type-theoretic approach of Refs. [4, 5].*

The proof is provided in App. D.3.

Therefore, working with projectors is not only a handy way to define the characterization constraints implied by type theory: There is a correspondence between the semantic definition of new types using $\{1, (,), \rightarrow\}$ and the construction of new projectors using $\{\neg, \otimes, \rightarrow\}$. If the base state structures are chosen to be quantum states, then the type theory of Ref. [4, 5] is recovered. In Sec. 6.2, we will provide an example in which the results of these papers are recovered from the projector algebra.

4 Implications of replacing types with projectors

We showed in the previous section that we can recast the type-theoretic characterization of higher-order quantum transformations into a projective formulation. In this section, we present the two advantages of doing so. The first one is the flexibility when defining the base types, whereas the second one is a new tool for comparing any two

types. While the first advantage is ‘only’ a formalisation of something that could have been done using the type-theoretic approach, the second one is intrinsic to working with projectors. It further reveals that the algebraic rules obeyed by the projectors are not arbitrary. They actually form an algebra, as well as a model of logic (first considered in Ref. [9]). This, on the one hand, provides a mathematical reason why the type system works the way it does. Moreover, it offers a new way of bridging between the type-theoretical approach [4, 5] and the categorical one [9, 26]. On the other hand, it allows for a classification of higher-order quantum processes, which will be shown to be based on the notion of signaling in Sec. 5.4, after a new way of combining projectors is introduced in Prop. 4.

4.1 Extending the hierarchy

Type theory is dependent on a set of base types. In Ref. [5], the first nontrivial types in the hierarchy, called the *elementary types*, are taken as the set of quantum states as in Eqs. (8). In terms of the projective characterization, this means that each base type is characterized with projector $\mathcal{I}$, and the elementary type on k subsystems is the set of quantum states characterized by projector $\mathcal{I}_A \otimes \mathcal{I}_B \otimes \dots \otimes \mathcal{I}_K$. The other types of higher-order transformations are subsequently obtained by using different sets of connectors to combine the k base projectors.

Within the framework of state structures, instead of having a fixed base type, one can consider any state structure as a base. In fact, any set of higher-order transformations can be coarse-grained into a base state structure by seeing all the systems they act on as a single overall system. With such a procedure, one can obtain many inequivalent base state structures, so the formalism can be further broadened to allow the parties to choose between several base state structures. For example, Alice could be choosing between inputting a quantum or a classical state, and the formalism would accommodate it because these two choices are valid state structures whose projectors commute with one another. In terms of projectors, it means that, for a given subsystem, the base projector can be picked from a set of pairwise commuting projectors on operator systems $\{\mathcal{P}, \mathcal{P}', \dots\}$, instead of being $\mathcal{I}$. Indeed, as shown in the following, building multipartite pro-

jectors from any set of commuting base projectors will lead to a set of commuting multipartite projectors. Thus, one can always coarse-grain these new multipartite projectors into a new set of base projectors, build new higher-order objects using them, and then repeat the operation to obtain a recursively more complex higher-order transformation (an example of such a procedure is presented in App. G).

Yet, and this is a difference with the type framework, the base projectors do not have to be built recursively by starting with $\mathcal{I}$. Combining several such projectors using $\{\neg, \otimes, \rightarrow\}$ to obtain a new set of multipartite projectors and coarse-graining into a new set of base projectors is not the only way of obtaining base projectors that are different from $\mathcal{I}$. Actually, any set of pairwise commuting projectors obeying Def. 2 can be considered. Therefore, different higher-order hierarchies than the one of quantum processes can be considered in principle; one only needs to start from a base projector characterizing a different set than the one of density matrices. For example, the projector to diagonal matrices in a fixed basis, corresponding to the state structure of classical states we mentioned above, is an example of a projector that cannot be obtained from $\mathcal{I}$ using $\{\neg, \otimes, \rightarrow\}$. We will come back to this at the end of the section, leading to the example presented in App. F. It is also discussed in more depth in App. A.1. Still, we emphasize that the base projectors one can consider are limited to sets of pairwise commuting elements.

4.2 The algebra of projectors

For any hierarchy of higher-order transformations, a key question is how to classify the non-equivalent kinds of transformations, and how to compare them. For example, it is known that the quantum supermaps happen to be a special kind of bipartite quantum channels [3, 4], but is there a systematic way to infer such a conclusion?

This is where characterizing using projectors helps: instead of focusing on types, the comparison can be made at the level of the linear subspaces defined by the projectors. Comparing now becomes straightforward as it amounts to computing the *overlap* between two subspaces. Moreover, this overlap does not require to be computed explicitly; the projector characterizing it – which is a function of the projectors associated with the

compared subspaces – is enough. Hence the comparison of two types, simplified into finding the overlap of their associated subspaces, can be further simplified into algebraic operations on their associated projectors. For any two projectors on operator systems $\mathcal{P}$ and $\mathcal{P}'$, we define the operation $\cap$ (nicknamed ‘cap’) as a bilinear function resulting in a new projector $\mathcal{P} \cap \mathcal{P}'$ called their *intersection* and characterizing the overlap of their two associated subspaces. Obviously, this is also a projector on an operator system.

4.2.1 Intersection, union, and lattice of projectors

If the projectors $\mathcal{P}$ and $\mathcal{P}'$ are commuting, the intersection is simply their composition:

$$\mathcal{P} \cap \mathcal{P}' \equiv \mathcal{P} \circ \mathcal{P}', \quad (25)$$

and it commutes with both $\mathcal{P}$ and $\mathcal{P}'$. Similarly, the following will be called the *union* of $\mathcal{P}$ and $\mathcal{P}'$:

$$\mathcal{P} \cup \mathcal{P}' \equiv \mathcal{P} + \mathcal{P}' - \mathcal{P} \cap \mathcal{P}'. \quad (26)$$

It characterizes the linear span of their two subspaces and the connector $\cup$ will be nicknamed ‘cup’. (See App. A.2 for the mathematical details of $\cap$ and $\cup$.)

With these, one can prove that an operator system characterized by $\mathcal{P}'$ is contained within another characterized by $\mathcal{P}$ by showing either of the following:

$$\mathcal{P} \cap \mathcal{P}' = \mathcal{P}'; \quad (27a)$$

$$\mathcal{P} \cup \mathcal{P}' = \mathcal{P}. \quad (27b)$$

In terms of projectors, these conditions will be concisely noted

$$\mathcal{P}' \subseteq \mathcal{P}. \quad (28)$$

In the following, any such equation comparing two projectors (i.e., inclusions and equivalences) should be understood as a statement about their images. When comparing state structures, the inclusion of projectors is sufficient to show inclusion (up to normalization) since the other constraint, positivity, is common to all state structures.

As it turns out, the projectors form an algebra over $\mathbb{C}$ under the connectives $\{\cap, \cup\}$ (see App. A.2 for the proof). And because every operator system must contain the identity, every projector in the algebra is contained between the depolarizing and identity projectors,

$$\mathcal{D} \subseteq \mathcal{P} \subseteq \mathcal{I}. \quad (29)$$

Actually, this is but the statement that a set of commuting projectors closed under a ‘join’ and ‘meet’ operations (intersection and union, respectively) forms a specific kind of algebra called a bounded lattice. Conditions (27) define the partial order (28), and Eq. (29) states that $\mathcal{D}$ and $\mathcal{I}$ are respectively the least and greatest elements in the lattice. Hence, the projectors can be compared using the partial order, which corresponds to the embedding of an operator system into another one. In terms of state structures, Eqs. (27) means that any state structure contains or is contained in another one (they can also both contain and be contained in another one, in which case the state structure are equivalent up to normalization); Eq. (28) then implies that the state structures can be classified using this inclusion relation as it defines a partial ordering; and Eq. (29) implies that the existence of a common least and greatest element in the ordering: every state structure is a subset of the set of density matrices (up to normalization) and a superset of the single-element set $\{\mathbb{1}\}$. In addition, it also implies that new state structures can be found by taking the intersections and unions of those that are known, and that the number of new state structures found using this procedure is finite.

Nevertheless, compared to $\otimes$ and $\neg$, the cap and the cup are new rules that cannot be reexpressed using the $\rightarrow$ connector alone. For example, the projector $\mathcal{I}_A$ can be obtained from $\mathcal{P}_A$ by using $\{\rightarrow, 1, (,), \cap, \cup\}$ since $\mathcal{I}_A = \mathcal{P}_A \cup \overline{\mathcal{P}}_A$ (see below), whereas by using $\{\rightarrow, 1, (,)\}$ one can only obtain $\mathcal{P}_A$ and $\overline{\mathcal{P}}_A$. So the downside of the algebra is that it goes outside of the type-theoretic framework of Ref. [4, 5].

Adding $\neg$ as an operation in the algebra promotes it into a *Boolean algebra* (or *Boolean lattice*), i.e. an algebra of idempotent elements which possess a *negation*. $\neg$ acts as the negation since it makes $\overline{\mathcal{P}}$ complementary to $\mathcal{P}$, $\mathcal{P} \cup \overline{\mathcal{P}} = \mathcal{I}$, and since it is an involution, $\overline{\overline{\mathcal{P}}} = \mathcal{P}$. A Boolean algebra obeys the De Morgan law:

$$\overline{\mathcal{P} \cap \mathcal{P}'} = \overline{\mathcal{P}} \cup \overline{\mathcal{P}'}, \quad (30)$$

which induces the following duality:

$$\mathcal{P}' \subseteq \mathcal{P} \iff \overline{\mathcal{P}} \subseteq \overline{\mathcal{P}'}, \quad (31)$$

(see App. A.3 for the proofs). These two rules greatly reduce the number of computations required to determine the relations between state

structures, as inclusions between functionals are dual to inclusions between the corresponding states.

4.2.2 The sub-lattice over a no signaling subset

The Boolean algebra of projectors is a lattice, meaning that every projector is either the union or the intersection of some other projectors. One can then characterize every other state structure related to the one(s) under consideration by enumerating all their possible intersections, unions, and negations. However, such an approach only considers each state structure as a global thing, characterized by a global projector $\mathcal{P}$, resulting in a simplistic lattice $\{\mathcal{P}, \overline{\mathcal{P}}, \mathcal{I}, \mathcal{D}\}$. But this is ignoring the finer structure of the projectors: often, they represent state structures defined on several subsystems, in which case they split as some combination of the projectors associated to each subsystem, say $\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K$, linked together to form $\mathcal{P}$ by using any mixture of the rules $\{\overline{\cdot}, \otimes, \rightarrow\}$, say $\mathcal{P} = (\mathcal{P}_A \otimes \overline{\mathcal{P}_B} \otimes \dots) \rightarrow \mathcal{P}_K$ for instance. It is then interesting to understand how this finer structure influences the overall shape of the lattice it generates: given a set $\{\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K\}$ of k base projectors, what can we say about the lattice of projectors acting on $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B \otimes \dots \otimes \mathcal{H}^K)$ obtained from every possible way to combine of these k projectors under $\{\overline{\cdot}, \otimes, \rightarrow\}$?

Before we begin, note that the assumption that all such projectors commute with each other becomes justified by the following lemma:

Lemma 2. *The rules $\{\overline{\cdot}, \otimes, \rightarrow\}$ preserve commutation.*

This is proven by direct computation, and so are all the relations presented in this subsection. (These proofs are all gathered in App. A.3 and A.4.)

Note also that negating the input of a transformation like $\mathcal{A} \rightarrow \mathcal{B}$ allows one to treat it as a composition akin to $\mathcal{A} \otimes \mathcal{B}$. Simply put, if $\mathcal{A} \rightarrow \mathcal{B}$ is a transformation from an input in $\mathcal{L}(\mathcal{H}^A)$ to an output in $\mathcal{L}(\mathcal{H}^B)$ (diagrammatically, it has an input wire at A and an output wire at B), negating the input turns it into a bipartite state preparation since it now has an output in $\mathcal{L}(\mathcal{H}^A)$ and one in $\mathcal{L}(\mathcal{H}^B)$ (diagrammatically, it now has an output wire at A as well as one at B).

Indeed, while an element of $\mathcal{A} \rightarrow \mathcal{B}$ can be seen as some composition of an effect at A (its input) together with a state at B (its output), an element of $\overline{\mathcal{A}} \rightarrow \mathcal{B}$ is accordingly the composition of the dual of an effect, i.e., a state, together with a state, yielding a bipartite state (with outputs at A and B). This composition is however different than the no signaling one $\otimes$. Thus, there are two different ways of composing the state structures $\mathcal{A}$ and $\mathcal{B}$ into a bipartite state structure: $\mathcal{A} \otimes \mathcal{B}$ and $\overline{\mathcal{A}} \rightarrow \mathcal{B}$. Under the same logic, taking the dual (negation) of one of the state structures in a no signaling composition, like $\overline{\mathcal{A}} \otimes \mathcal{B}$ for instance, allows to treat it as a transformation: it has an input (effect of $\overline{\mathcal{A}}$) and an output (state of $\mathcal{B}$), so it can be interpreted as a map from a valid state in $\mathcal{A}$ to one in $\mathcal{B}$.

In terms of projectors, the operation $\overline{\cdot} \rightarrow \cdot = \overline{\overline{\cdot} \otimes \overline{\cdot}}$ is then the ‘compositional’ version of the transformation $\rightarrow$,

$$\overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B = \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}}. \quad (32)$$

It is a ‘larger’ composition than the tensor (see App. A.4.1 for the proof) as $\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \neq \mathcal{P}_A \otimes \mathcal{P}_B$ and

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}}. \quad (33)$$

Moreover, like the tensor, this $\overline{\cdot} \rightarrow \cdot$ composition is associative since

$$\overline{(\overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B)} \rightarrow \mathcal{P}_C = \overline{\mathcal{P}_A} \rightarrow (\overline{\mathcal{P}_B} \rightarrow \mathcal{P}_C). \quad (34)$$

Indeed, the left-hand side of the above corresponds to the partition $(AB)C$, i.e. first applying $\overline{\cdot} \rightarrow \cdot$ between A and B , then between (AB) and C ; whereas the right-hand side corresponds to the partition $A(BC)$. This is readily proven by expressing $\cdot \rightarrow \cdot$ as $\overline{\cdot \otimes \overline{\cdot}}$ in Eq. (34): the left-hand side becomes $\overline{(\overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B)} \rightarrow \mathcal{P}_C = \overline{(\overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}})} \rightarrow \mathcal{P}_C = \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \otimes \overline{\mathcal{P}_C}}$ and, similarly, the right-hand side becomes $\overline{\mathcal{P}_A} \rightarrow (\overline{\mathcal{P}_B} \rightarrow \mathcal{P}_C) = \overline{\mathcal{P}_A} \rightarrow (\overline{\overline{\mathcal{P}_B} \otimes \overline{\mathcal{P}_C}}) = \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \otimes \overline{\mathcal{P}_C}}$. This proof is direct to generalize by induction: no matter the ordering of parentheses and the number of parties, one always ends up with the same expression when expressed using

the negation and the tensor, e.g.

$$\begin{aligned}
\overline{\mathcal{P}_A} &\rightarrow (\overline{\mathcal{P}_B} \rightarrow \dots (\overline{\mathcal{P}_J} \rightarrow \mathcal{P}_K) \dots) \\
&= \overline{(\mathcal{P}_A \rightarrow \mathcal{P}_B)} \rightarrow (\dots \rightarrow (\overline{\mathcal{P}_J} \rightarrow \mathcal{P}_K) \dots) \\
&= \dots \\
&= \overline{(\dots (\overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B) \rightarrow \dots) \rightarrow \mathcal{P}_J} \rightarrow \mathcal{P}_K \\
&= \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \otimes \dots \overline{\mathcal{P}_J} \otimes \overline{\mathcal{P}_K}}. \quad (35)
\end{aligned}$$

Going back to the comparison of higher-order transformations, we have so far argued that it is reduced to the study of lattices of commuting projectors on operator systems. For k parties, any expression $\mathcal{P}$ built using $\{\overline{\cdot}, \otimes, \rightarrow\}$ is contained in a Boolean lattice closed under operations $\{\cap, \cup\}$ with smallest element $\mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K$ and biggest element $\mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K$ in accordance with Eq. (29). Of course, with an increasing number of parties these lattices quickly get huge. However, any projector in the lattice has natural subset and superset that are in general different than the tensor product of the depolarizing and of the identity projectors, meaning that projectors usually belong to a sublattice. These sets are built using exclusively one of the two possible combination rules we just discussed.

Definition 7 (No signaling subset / Fully signaling superset). *Let $\mathcal{P}$ be a projector on an operator system built using k base projectors $\mathcal{P}_A, \mathcal{P}_B \dots \mathcal{P}_K$ composed together using operations $\{\overline{\cdot}, \cap, \cup, \otimes, \rightarrow\}$.*

*The projector $\mathcal{P}_{NS}$ to its **no signaling subset** is the largest projector contained in $\mathcal{P}$ obtained as the tensor composition of single-partite projectors. By construction, the single-partite projector associated with party X can only be the base projector, its negation, or the depolarizing superoperator $\mathcal{D}_X$ so that $\mathcal{P}_{NS}$ reads*

$$\mathcal{P}_{NS} \equiv \tilde{\mathcal{P}}_A^{\mathcal{D}} \otimes \tilde{\mathcal{P}}_B^{\mathcal{D}} \otimes \dots \otimes \tilde{\mathcal{P}}_K^{\mathcal{D}} \subseteq \mathcal{P}, \quad (36)$$

where the $\tilde{\mathcal{P}}_X^{\mathcal{D}}$ notation means an element chosen among $\{\mathcal{P}_X, \overline{\mathcal{P}_X}, \mathcal{D}_X\}$ depending on the exact form of $\mathcal{P}$.

*The projector $\mathcal{P}_{FS}$ to its **fully signaling superset** is the smallest projector containing $\mathcal{P}$ obtained as the “ $\overline{\cdot} \rightarrow \cdot$ ” composition of single-partite projectors. By construction, the single-partite projector associated with party X can only*

be the base projector, its negation, or the identity superoperator $\mathcal{I}_X$ so that $\mathcal{P}_{FS}$ reads

$$\mathcal{P} \subseteq \overline{\tilde{\mathcal{P}}_A^{\mathcal{I}}} \otimes \overline{\tilde{\mathcal{P}}_B^{\mathcal{I}}} \otimes \dots \otimes \overline{\tilde{\mathcal{P}}_J^{\mathcal{I}}} \otimes \overline{\tilde{\mathcal{P}}_K^{\mathcal{I}}} \equiv \mathcal{P}_{FS}, \quad (37)$$

where the $\tilde{\mathcal{P}}_X^{\mathcal{I}}$ notation means an element chosen among $\{\mathcal{P}_X, \overline{\mathcal{P}_X}, \mathcal{I}_X\}$ depending on the exact form of $\mathcal{P}$.

The reason behind the names of these sets will be explained in the next section. To see why these sets always exist, notice that the definition allows $\mathcal{P}_{NS}$ to be equal to $\mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K$ and $\mathcal{P}_{FS}$ to be $\mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K$ in the worst-case scenario. In general, these define a sub-lattice

$$\mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K \subseteq \mathcal{P}_{NS} \subseteq \mathcal{P} \subseteq \mathcal{P}_{FS} \subseteq \mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K, \quad (38)$$

around each projector $\mathcal{P}$.

The point is that a projector $\mathcal{P}$ is always embedded in a lattice of projectors that use the same base projectors and is contained between $\mathcal{I}$ and $\mathcal{D}$. But it is often uninteresting to compare this projector to elements that have a different interpretation. In the single-partite case for instance, the lattice is $\{\mathcal{P}, \overline{\mathcal{P}}, \mathcal{I}, \mathcal{D}\}$, but it is not very enlightening to compare $\mathcal{P}$ with $\overline{\mathcal{P}}$ as one is the dual of the other, meaning that while one is related to state preparation, the other is related to measurements (effects), which are qualitatively different. The following lemma tells us that a projector $\mathcal{P}$ is also embedded in a smaller lattice of only the projectors that have the same interpretation (in terms of seeing a subsystem as either state or effect) as $\mathcal{P}$.

Lemma 3 (No signaling sub-lattice around a projector). *Let $\mathcal{P}$ be a projector on an operator system built using k base projectors $\mathcal{P}_A, \mathcal{P}_B \dots \mathcal{P}_K$ composed together using operations $\{\overline{\cdot}, \otimes, \rightarrow\}$. Then, it belongs to a sub-lattice closed under operations $\{\cap, \cup\}$ so that*

$$\tilde{\mathcal{P}}_A \otimes \dots \otimes \tilde{\mathcal{P}}_K \subseteq \mathcal{P} \subseteq \overline{\overline{\tilde{\mathcal{P}}_A}} \otimes \dots \otimes \overline{\overline{\tilde{\mathcal{P}}_K}}, \quad (39)$$

where the $\tilde{\mathcal{P}}_X$ notation means an element chosen among $\{\mathcal{P}_X, \overline{\mathcal{P}_X}\}$ depending on $\mathcal{P}$. Importantly, the choice is the same on both sides of the equation.

This follows mainly from the property (136) that $\mathcal{P}_A \subseteq \mathcal{P}'_A \subseteq \mathcal{P}''_A \Rightarrow \mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}'_A \otimes \mathcal{P}_B \subseteq$

$\mathcal{P}_A'' \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}_A'' \otimes \mathcal{P}_B}$, and from property (31). The complete proof is provided in App. D.4.

A rule of thumb for finding the no signaling projectors is to ‘count the number of bars’ above each base projector in a given expression $\mathcal{P}$. Since negation is an involution, an odd (respectively, even) amount will indicate that the projector is (not) negated in the no signaling subset. The fully signaling projector is then obtained by taking the negation of the no signaling one. More rigorously, the projectors of Lemma 3 are found by repetitively using Eq. (33). For example, to find the no signaling subset of $(\mathcal{P}_A \otimes (\mathcal{P}_B \rightarrow \mathcal{P}_C)) \rightarrow \mathcal{P}_D$, one first expresses it using negations and tensors products, $(\mathcal{P}_A \otimes (\mathcal{P}_B \rightarrow \mathcal{P}_C)) \rightarrow \mathcal{P}_D = \mathcal{P}_A \otimes \mathcal{P}_B \otimes \overline{\mathcal{P}_C} \otimes \overline{\mathcal{P}_D}$, then by inspection $\mathcal{P}_A$ and $\mathcal{P}_C$ have an odd number of negations above them (one and three, respectively) so they will be negated, whereas $\mathcal{P}_B$ and $\mathcal{P}_D$ have an even number (two and two) so they will not be. On that account, the no signaling subset of $(\mathcal{P}_A \otimes (\mathcal{P}_B \rightarrow \mathcal{P}_C)) \rightarrow \mathcal{P}_D$ is $\overline{\mathcal{P}_A} \otimes \mathcal{P}_B \otimes \overline{\mathcal{P}_C} \otimes \mathcal{P}_D$ and its fully signaling superset is $\mathcal{P}_A \otimes \overline{\mathcal{P}_B} \otimes \mathcal{P}_C \otimes \overline{\mathcal{P}_D}$.

Thus, for a given $\mathcal{P}$, the state structure it defines is always contained between a no signaling one and a fully signaling one. Note that the position of the negations in Eqs. (36) and (37) are the same because of Eq. (33). Lemma 3 indeed states that the projectors to these subspaces are actually the least and greatest elements of a sub-algebra stable under $\cap$ and $\cup$, i.e. a sub-lattice. A rule of thumb to determine if a projector is in the no signaling sub-lattice of another one is again to ‘count the number of bars’ above each base projector. If each number is the same modulo two, then it belongs to it, otherwise it does not. For example, the no-signaling lattice around projector $\mathcal{P}_A$ is $\{\mathcal{P}_A\}$ itself (in which case it is both its no signaling sub- and fully signaling super-sets), whereas the one around $\mathcal{P}_A \otimes \mathcal{P}_B$ is $\{\mathcal{P}_A \otimes \mathcal{P}_B, \overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B\}$ because in $\overline{\mathcal{P}_A} \rightarrow \mathcal{P}_B = \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}$ there are ‘two bars’ over both $\mathcal{P}_A$ and $\mathcal{P}_B$ while in $\mathcal{P}_A \otimes \mathcal{P}_B$ there are ‘zero bars’ above the two.

Lemma 3 is then a ‘decompositional’ approach to the no signaling sub-lattice: given a multipartite projector, one considers every other projector in its sub-lattice. This sub-lattice w.r.t. a given $\mathcal{P}$ can be obtained by defining the projectors $\tilde{\mathcal{P}}_X$ involved in its no signaling subset as a

new set of base projectors, $\tilde{\mathcal{P}}_X \mapsto \mathcal{P}_X$ (i.e., by redefining the base projectors so to incorporate the negations), and by restricting the choice of operations $\{\neg, \otimes, \rightarrow\}$ to a choice that preserves the same no signaling subset, which, from of Eq. (33), is $\{\cdot \otimes \cdot, \neg \rightarrow \cdot\}$ (this restriction is the formal version of the rule of thumb mentioned above, since it restricts the operations to only those ‘adding an even number of bars over each projector’). The procedure we just described is then the ‘compositional’ approach to the no signaling sub-lattice: given a set of base projectors, one constructs the no signaling sub-lattice out of it. This leads to the following counterpart to Lemma 3.

Lemma 4 (Building a no signaling sub-lattice.). *Let $\{\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K\}$ be a set of k projectors on operator systems each associated with a specific Hilbert space $\mathcal{H}^A, \mathcal{H}^B, \dots, \mathcal{H}^K$, i.e., a set of base projectors. Any expression $\mathcal{P}$ built using each element of this set once and under the rules $\{\cdot \otimes \cdot, \neg \rightarrow \cdot\}$ is a projector on an operator system over $\mathcal{L}(\mathcal{H}^A \otimes \dots \otimes \mathcal{H}^K)$ and belongs to the no signaling lattice*

$$\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K \subseteq \mathcal{P} \subseteq \overline{\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K}. \quad (40)$$

This will be proven for a more general setting in Proposition 6, but we first need to discuss the notion of no signaling to better understand why these sub-lattices are called ‘no signaling’. This will be the topic of the next section.

If needed, concrete examples of the use of the algebra of projectors developed in this section are provided in App. E. As we mentioned, the algebra is well-defined for any base projector commuting with the identity and depolarizing projectors. Using base projectors different from ones that can be obtained from $\mathcal{I}$ and the operations of the algebra amounts to building a theory that is different than quantum theory and its higher-order generalizations. An example of such a toy model, which we call ‘biased quantum theory’, is presented in App. F. These two appendices come with many remarks providing intuition on the notion of no signaling at the level of state structures so to motivate the next section.

Remark: the algebra of projectors is (almost) Linear Logic. With the further additions of the tensor and of the transformation operations, the Boolean algebra of projectors is

lifted to an abstract lattice-like structure. Such a structure is always associated with a model of formal logic. Observe that the transformation between states is equivalent to the reverse transformation between effects,

$$\mathcal{P}_A \rightarrow \mathcal{P}_B = \overline{\mathcal{P}}_A \leftarrow \overline{\mathcal{P}}_B = \overline{\mathcal{P}_A \otimes \overline{\mathcal{P}}_B}. \quad (41)$$

This is a sign of the algebra being an instance of logic: it follows the logic rule that if A implies B, then not B must imply not A. ‘Propositions’ $\mathcal{P}$ consist of k ‘sub-propositions’ $\{\mathcal{P}_A, \dots, \mathcal{P}_K\}$ that can be composed in two different manners: using $\otimes$ or $\overline{\cdot} \rightarrow \cdot$; and each (sub)proposition can be negated using $\overline{\cdot}$. To compare two propositions, the connectives $\cap$ and $\cup$ are used, and the comparison of two propositions is itself a valid proposition. The lattice is then the set of all propositions. As a model of logic, the projector algebra happens to form a degenerate instance of *multiplicative additive linear logic* (MALL) [14] because its connectives obey the following De Morgan dualities:

$$\overline{\overline{\mathcal{P}}_A} = \mathcal{P}_A; \quad (42a)$$

$$\overline{\mathcal{P}_A \cap \mathcal{P}'_A} = \overline{\mathcal{P}}_A \cup \overline{\mathcal{P}'_A}; \quad (42b)$$

$$\overline{\mathcal{P}_A \otimes \mathcal{P}_B} = \overline{\mathcal{P}}_A \rightarrow \overline{\mathcal{P}}_B. \quad (42c)$$

(See App. A.4.3 for more details.) This observation connects our formalism to the categorical approach of Ref. [9], in which they noticed that the logic of higher-order quantum transformations form an instance of *multiplicative linear logic* (MLL)¹¹. The interested reader is invited to consult the categorical counterpart of the present article [26], in which they build a non-degenerate MALL out of higher-order quantum theory and discuss the internal logic of the characterization in a rigorous way.

5 No signaling

In a previous work, the projective characterization was used to prove that the multi-round process matrix, an extension of the process matrix

¹¹In our language, the observation is that the projector rules $\{\overline{\cdot}, \cdot \otimes \cdot, \overline{\cdot} \rightarrow \cdot\}$ is an instance of MLL. The fact that the algebra of the projectors characterizing quantum higher-order transformations is an instance of MLL can additionally be understood as the consequence of the correspondence between the types of Ref. [5] and $\ast$ -autonomous categories as in Ref. [9].

[8], is a linear sum of quantum combs [11, Theorem 2]. By doing so, it relates an object that typically involves indefinite causal order (the multi-round process matrix, and as a consequence the process matrix) with *several* that do not (a quantum comb represents a causally ordered succession of quantum channels sharing a side channel; in the literature, this structure is usually referred to as a quantum network or as a quantum channel with memory), thus formalizing the intuitive idea that ‘an indefinite causal order is a superposition of *several* causal orderings’. Formally, the result states that the projector characterizing the MPM is exactly a sum of projectors characterizing quantum combs, with each term in the sum corresponding to each different causal ordering that the MPM can feature.

As the MPM and the quantum comb are typical instances of objects in structures that can be characterized using the projective characterization, this suggests that the projectors built using the previously defined operations can be split into terms that have a fixed signaling direction, so that the causal structure can be inferred from the analysis of the projector alone. Therefore, in this section, the notion of no signaling is extended to state structures in order for this decomposition to be systematized. Quasi-orthogonality will be shown to be the relevant notion for no signaling at the level of projectors, allowing in turn to define the *prec*, an algebraic rule for composing projectors that encodes a one-way signaling constraint.

5.1 Definition of no signaling for state structures

No signaling from the output to the input can be imposed on a transformation as an extra condition. This translates the physical idea that, if the output is assumed to be in the causal future of the input, no (deterministic) operation done in the future can influence, or signal to, the past [27].

The subset obeying this constraint will be noted using a “ $\prec$ ” symbol, nicknamed the **prec**. By convention, this connector will be seen as a composition, like $\otimes$, rather than a transformation, like $\rightarrow$. As explained in last section, this means that the two parts of the connector are treated on equal footing; an object of $\mathcal{A} \otimes \mathcal{B}$ or of $\mathcal{A} \prec \mathcal{B}$ is *composed* of objects in $\mathcal{A}$ and in $\mathcal{B}$. This implies that $\mathcal{A} \prec \mathcal{B}$ should contain $\mathcal{A} \otimes \mathcal{B}$

as its no signaling subset. On the contrary, an object of $\mathcal{A} \rightarrow \mathcal{B}$ transforms an input in $\mathcal{A}$ into an output in $\mathcal{B}$; it needs an object of $\mathcal{A}$, and therefore it is composed of objects in $\overline{\mathcal{A}}$ and in $\mathcal{B}$. This implies that $\mathcal{A} \rightarrow \mathcal{B}$ contains $\overline{\mathcal{A}} \otimes \mathcal{B}$ as its no signaling subspace instead. With respect to the sub-lattice of Lemma 4, we want to introduce a third connective rule, $\prec$, as an in-between $\otimes$ and $\rightarrow$. This rule will encode a *composition* that allows signaling from the state structure on its left to the state structure on its right. In order to define our use of the term signaling with respect to *transformations*, the left-hand side of connectives $\otimes$ and $\prec$ will therefore be negated; we are looking for the particular subset $\overline{\mathcal{A}} \prec \mathcal{B}$ of the transformation $\mathcal{A} \rightarrow \mathcal{B}$ that forbids signaling from output to input. In other words, we want to induce a fine-graining of the order relation (33) into $\overline{\mathcal{A}} \otimes \mathcal{B} \subseteq \overline{\mathcal{A}} \prec \mathcal{B} \subseteq \mathcal{A} \rightarrow \mathcal{B}$.

The proposed condition for defining the set $\overline{\mathcal{A}} \prec \mathcal{B}$ is the requirement that the input of the transformation should be independent of the *particular choice* of effect applied at its output:

$$M \in \overline{\mathcal{A}} \prec \mathcal{B} \subseteq \mathcal{A} \rightarrow \mathcal{B} : \quad \forall N, N' \in \overline{\mathcal{B}}, \quad \text{Tr}_B [M \cdot (\mathbb{1} \otimes N)] = \text{Tr}_B [M \cdot (\mathbb{1} \otimes N')]. \quad (43)$$

As the identity matrix $\mathbb{1}$ is always both a valid state and effect (with suitable normalization) in this work, this requirement is rephrased as the following constraint,

$$M \in \overline{\mathcal{A}} \prec \mathcal{B} : \quad \forall N \in \overline{\mathcal{B}}, \quad \text{Tr}_B [M \cdot (\mathbb{1} \otimes N)] = \text{Tr}_B \left[M \cdot \left(\mathbb{1} \otimes \frac{\mathbb{1}}{c_B} \right) \right]. \quad (44)$$

This is depicted in Fig. 6b. When the above is satisfied, we will use it as a definition to say that M is *no signaling* from B to A [17, 18]. In the same way, one can define the no input-to-output signaling subset, $\overline{\mathcal{A}} \succ \mathcal{B}$, by requiring the reverse condition, as drawn in Fig. 6f.

For example, if $M \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1} \otimes \mathcal{H}^{B_0} \otimes \mathcal{H}^{B_1})$ is a bipartite quantum channels, the local inputs on Alice's side has the form $\rho_{A_0} \otimes \mathbb{1}_{A_1} \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$ where ρ is a quantum state, and the reduced single-partite channel seen by Bob after Alice has acted is given by $\text{Tr}_{A_0 A_1} [M \cdot (\rho_{A_0} \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_0} \otimes \mathbb{1}_{B_1})]$ (see examples E.3 and E.4 if this is unclear). We then say that the channel M is no signaling

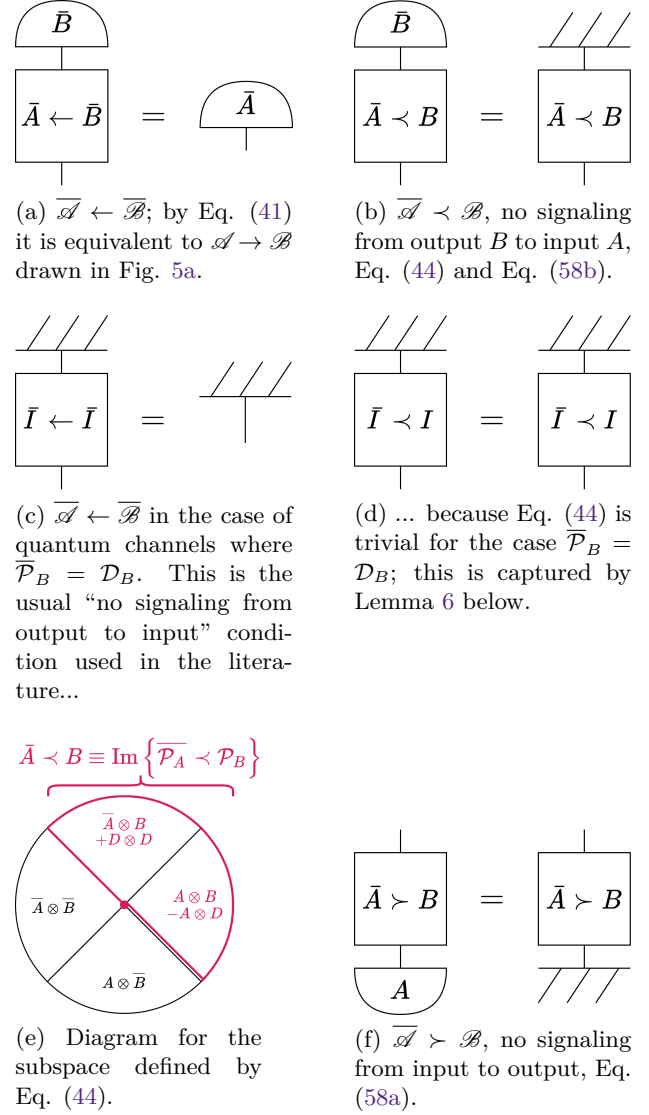


Figure 6: Graphical interpretation of the defining condition of a transformation and of its no signaling from output to input subset. On the first line, the condition to be a valid element of $\mathcal{A} \rightarrow \mathcal{B}$ is represented as the equivalent condition $\overline{\mathcal{A}} \leftarrow \overline{\mathcal{B}}$ in 6a. The extra condition that the output cannot use the transformation to signal to the input, Eq. (44), is represented in 6b. On the second line, the same two conditions are represented in the case where $\overline{\mathcal{P}}_B = \mathcal{D}_B$. On the last line, the diagram representing the subspace defined by condition (58b) of Fig. 6b is drawn in Fig. 6e as an illustration of Prop. 8. Finally, the reverse condition, no signaling from input to output, Eq. (58a), has its graphical representation in Fig. 6f.

from Alice to Bob, if, no matter which input $\rho \otimes \mathbb{1}$ Alice chooses, Bob always sees the same channel on his side. In this case, she has a single choice of input at A_1 , $\mathbb{1}$, but she can choose any quantum state she wants at A_0 . The channel is

therefore no signaling from Alice to Bob if the following holds for any two states ρ, σ ,

$$\begin{aligned} \text{Tr}_{A_0 A_1} [M \cdot (\rho_{A_0} \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_0} \otimes \mathbb{1}_{B_1})] = \\ \text{Tr}_{A_0 A_1} [M \cdot (\sigma_{A_0} \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_0} \otimes \mathbb{1}_{B_1})] . \end{aligned} \quad (45)$$

Remark: relation with the previous definitions (see Refs. [17, 18, 3, 9]). In App. C, condition (43) is rephrased to correspond to a constraint on the probability distribution of the outcomes associated with the resolutions of some $N \in \overline{\mathcal{B}}$ acting on a bipartite state $\mathcal{A} \otimes \mathcal{B}$. This is done in order to link it to the usual theory-independent notion of no signaling used in nonlocality. With respect to the notion used in the case of quantum channel formalism (which gives Eq. (46) below), our notion of no signaling actually requires more conditions, but recovers the quantum channel notion as a special case. Indeed, observe that due to Eq. (41), a map $M \in \mathcal{A} \rightarrow \mathcal{B}$ can be equivalently interpreted as $\overline{\mathcal{A}} \leftarrow \overline{\mathcal{B}}$. Consequently, M must obey $\forall N_B \in \overline{\mathcal{B}}$, $\text{Tr}_B [M (\mathbb{1}_A \otimes N_B)] \in \overline{\mathcal{A}}$ by definition (see Fig. 6a), and since $\overline{\mathcal{A}} = \{\mathbb{1}_A\}$, we attain the usual $M \in \mathcal{A} \rightarrow \mathcal{B} \iff$

$$\text{Tr}_B [M (\mathbb{1}_A \otimes \mathbb{1}_B)] = \text{Tr}_B [M] = \mathbb{1}_A , \quad (46)$$

which is used in previous works [3, 9] as a condition to check if M is a valid quantum channel *but at the same time* as one to check if there is no signaling from output to input (depending on the source, this condition, Eq. (46), which is graphically depicted in Fig. 6c, is also called “causality condition” [9]). The reason why using Eq. (46) for the definition of the quantum channel, as well as the extra condition that the output is no signaling to the input, instead of using Eqs. (46) and (44), is that Eq. (44) becomes tautological in that case, so it is automatically satisfied (see Fig. 6d). However, this equivalence between a general transformation and a no signaling one is only valid for the case of density matrices (as proven explicitly in Lemma 6 below) and hides the subtle difference between the two definitions in the general case: being a valid transformation between state structures does not guarantee that the output cannot be used to signal to the input. Graphically, the condition represented in Fig. 6a) does not imply the one in Fig. 6b).

5.2 Projective characterization

Equation (44) can be rewritten as

$$\begin{aligned} \forall N \in \overline{\mathcal{B}}, \\ \text{Tr}_B [M \cdot (\mathbb{1} \otimes N)] = \frac{1}{d_B} \text{Tr}_B [M] \text{Tr} [N] . \end{aligned} \quad (47)$$

This ‘factorization’ of the partial trace is reminiscent of the property (17) of quasi-orthogonal subspaces. In App. C, it is proven that our definition of no signaling is indeed enforced by the following subsystem-wise notion of quasi-orthogonality.

Lemma 5. *Let M be an operator in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$, let V be one in $\mathcal{A} \subset \mathcal{L}(\mathcal{H}^A)$, and let N be one in $\mathcal{L}(\mathcal{H}^B)$. Then a necessary and sufficient condition for:*

$$\text{Tr}_A [M \cdot (V \otimes N)] = \frac{\text{Tr}_A [M] \text{Tr} [V]}{d_A} \cdot N , \quad (48)$$

to hold for all V and N is that

$$(\overline{\mathcal{P}}_A \otimes \mathcal{I}_B) \{M\} = M . \quad (49)$$

That is to say, that M belongs to $\overline{\mathcal{A}} \otimes \mathcal{L}(\mathcal{H}^B)$.

The proof is presented in App. D.5. The above lemma states that since no signaling condition on M —here encoded as Eq. (48)—is a linear constraint, there exists a subspace of all the operators obeying this conditions, and this subspace is the operator system $\overline{\mathcal{A}} \otimes \mathcal{L}(\mathcal{H}^B)$.

Now, the state structure $\overline{\mathcal{A}} \prec \mathcal{B}$ has been defined as the subset of the state structure $\mathcal{A} \rightarrow \mathcal{B}$ whose elements are no signaling from the output to the input. So, alongside the valid transformation constraint, encoded by the projector $\mathcal{P}_A \rightarrow \mathcal{P}_B$ (according to Prop. 2), the extra constraint (43) must be added, and it is encoded by the projector $\mathcal{I}_A \otimes \mathcal{P}_B$ (according to Lemma 5). As two projective constraints must be simultaneously satisfied, we can use the $\cap$ operation to combine them into a unique projector:

$$\begin{aligned} (\mathcal{P}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B) = \\ (\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \overline{\mathcal{P}}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B) \\ = \mathcal{I}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B . \end{aligned} \quad (50)$$

This projector, as it turns out, is a projector on operator system. Besides, neither positivity nor normalization has been affected by the addition of the new constraint. As a consequence, $\overline{\mathcal{A}} \prec \mathcal{B}$ is a state structure.

Proposition 4 (One-way signaling transformation). *Let $\mathcal{M} \in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^B))$ be a structure-preserving map between state structures $\mathcal{A}$ and $\mathcal{B}$ as in Def. 6. Let $\{M\} \equiv \overline{\mathcal{A}} \prec \mathcal{B}$ be the subset of all M in CJ representation that obey the no signaling from output to input constraint, Eq. (43), $\forall N, N' \in \overline{\mathcal{B}} : \text{Tr}_B[M \cdot (\mathbb{1} \otimes N)] = \text{Tr}_B[M \cdot (\mathbb{1} \otimes N')]$. Then, $\overline{\mathcal{A}} \prec \mathcal{B}$ is characterized by*

$$M \in \overline{\mathcal{A}} \prec \mathcal{B} \iff M \geq 0, \quad (51a)$$

$$\text{Tr}[M] = c_{\overline{\mathcal{A}}}c_B, \quad (51b)$$

$$(\mathcal{I}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B)\{M\} = M, \quad (51c)$$

and it is a state structure as in Def. 3 since

$$\mathcal{I}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \equiv \overline{\mathcal{P}}_A \prec \mathcal{P}_B \quad (52)$$

is a projector on an operator system as in Def. 2.

Proof. Positivity and trace conditions are inherited from M being a valid transformation. The projector condition is obtained directly from the above discussion, i.e. by taking the intersection of the projector on valid transformations, which is $\mathcal{P}_A \rightarrow \mathcal{P}_B$ by Prop. 2, with the projector on the set of operators obeying the no signaling condition, which is $\mathcal{I}_A \otimes \mathcal{P}_B$ by Lem. 5. That (52) is a projector on an operator system comes from the fact that it is obtained from two valid projectors on operator systems, $\mathcal{P}_A \rightarrow \mathcal{P}_B$ and $\mathcal{I}_A \otimes \mathcal{P}_B$, composed using the cap operation, $\cap$, which preserves this property (this is proven in App. A.5). $\square$

See Fig. 6e for a diagrammatic depiction of the subspace associated with projector (51c).

As is the case for the transformation, the one-way signaling transformation can be seen as the composition of the state structures $\overline{\mathcal{A}}$ and $\mathcal{B}$ whose characteristics are encoded in a specific algebraic operation that combines projectors $\overline{\mathcal{P}}_A$ and $\mathcal{P}_B$. Formally,

Definition 8 (One-way signaling composition). *Let $\mathcal{A}$ and $\mathcal{B}$ be two state structures as in Eqs. (7), their A-to-B one-way signaling composition is the state structure $\mathcal{A} \prec \mathcal{B} \subset \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$*

consisting of all operators W characterized by the following conditions:

$$W \geq 0, \quad (53a)$$

$$\text{Tr}[W] = c_A c_B, \quad (53b)$$

$$(\mathcal{P}_A \prec \mathcal{P}_B)\{W\} = W. \quad (53c)$$

where

$$\mathcal{P}_A \prec \mathcal{P}_B \equiv \mathcal{I}_A \otimes \mathcal{P}_B - \overline{\mathcal{P}}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (54)$$

The same way, their B-to-A one-way signaling composition is the analogously defined state structure $\mathcal{A} \succ \mathcal{B} \subset \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ but instead using the projector

$$\mathcal{P}_A \succ \mathcal{P}_B \equiv \mathcal{P}_A \otimes \mathcal{I}_B - \mathcal{D}_A \otimes \overline{\mathcal{P}}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (55)$$

5.3 Relating the compositions

The main rules for the characterization of state structures through their projector rules that have been defined so far are summarized in Table 1. Using the new connectives introduced in Sec. 4, i.e. $\cap$ and $\cup$ (defined by Eqs. (25) and (26), respectively), the three bipartite connectors that have been introduced so far, i.e. $\otimes$, $\prec$, and $\rightarrow$ (defined by Eqs. (19), (54)) and (24)) can be related together in the algebra of projectors.

For instance, consider the four ways of composing a state in $\mathcal{A}$ with an effect in $\mathcal{B}$. They lead to four projectors all having $\overline{\mathcal{P}}_A \otimes \mathcal{P}_B$ as their no signaling subset and the following relations can be inferred:

$$\overline{\mathcal{P}}_A \otimes \mathcal{P}_B = (\overline{\mathcal{P}}_A \prec \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B), \quad (56a)$$

$$\mathcal{P}_A \rightarrow \mathcal{P}_B = (\overline{\mathcal{P}}_A \prec \mathcal{P}_B) \cup (\overline{\mathcal{P}}_A \succ \mathcal{P}_B). \quad (56b)$$

See App. A.5.2 for the proof.

Using that $\prec$ represents one-way signaling composition, these have a concrete physical interpretation. Consider the set $\overline{\mathcal{A}} \otimes \mathcal{B}$ first. Its elements are transformations from $\mathcal{A}$ to $\mathcal{B}$ as each $M \in \overline{\mathcal{A}} \otimes \mathcal{B}$ satisfies $\forall V \in \mathcal{A}$:

$$\text{Tr}_A[M \cdot (V \otimes \mathbb{1})] \in \mathcal{B}. \quad (57)$$

This is because its linear span is included in the one of the transformations; projector-wise, this is the relation $\overline{\mathcal{P}}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B$. As the tensor product lives in a subspace of the transformation, each linear constraint obeyed by the latter must be obeyed by the former as well, so this includes Eq. (57). In addition to that, the right-hand side of Eq. (56a) puts two extra conditions similar to Eq. (44) on $\overline{\mathcal{A}} \otimes \mathcal{B}$ (see Figs. 6f and 6b):

Name	Characterization	Proj. rule	Rule nickname
State	Definition 3; Eqs. (7)	$\mathcal{P}_A$	/
Effect	Proposition 1; Eqs. (15)	$\overline{\mathcal{P}}_A$	Negation
2-ways Sign.	Proposition 2; Eqs. (23)	$\mathcal{P}_A \rightarrow \mathcal{P}_B$	Transformation
1-way Sign.	Proposition 4; Eqs. (51)	$\overline{\mathcal{P}}_A \prec \mathcal{P}_B$	Prec
No Sign.	Proposition 5; Eqs. (59)	$\overline{\mathcal{P}}_A \otimes \mathcal{P}_B$	Tensor

Table 1: Summary of the characterization rules for state structure of states, effects, and the three ways to transform from state structure $\mathcal{A}$ to $\mathcal{B}$.

Proposition 5 (No signaling transformation). *Let $\mathcal{M} \in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^B))$ be a structure-preserving map between state structures $\mathcal{A}$ and $\mathcal{B}$ as in Def. 6. Let $\{M\} \equiv \overline{\mathcal{A}} \otimes \mathcal{B}$ be the subset of all $\mathcal{M}$ in CJ representation that obeys the no signaling from input to output as well as no signaling from output to input constraints: $\forall V \in \mathcal{A}, \forall N \in \overline{\mathcal{B}},$*

$$\text{Tr}_A [M \cdot (V \otimes \mathbb{1})] = \text{Tr}_A \left[M \cdot \left(\frac{\mathbb{1}}{c_A} \otimes \mathbb{1} \right) \right]; \quad (58a)$$

$$\text{Tr}_B [M \cdot (\mathbb{1} \otimes N)] = \text{Tr}_B \left[M \cdot \left(\mathbb{1} \otimes \frac{\mathbb{1}}{c_B} \right) \right]. \quad (58b)$$

Then, $\overline{\mathcal{A}} \otimes \mathcal{B}$ is a state structure characterized by

$$M \in \overline{\mathcal{A}} \otimes \mathcal{B} \iff M \geq 0, \quad (59a)$$

$$\text{Tr}[M] = c_A c_B, \quad (59b)$$

$$(\overline{\mathcal{P}}_A \otimes \mathcal{P}_B) \{M\} = M. \quad (59c)$$

Proof. The only difference with Prop. 4 is the projector condition, $(\overline{\mathcal{P}}_A \otimes \mathcal{P}_B) \{M\} = M$. By using twice Lemma 5, Eqs. (59) constraint the subspace to the intersection $(\mathcal{P}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \otimes \mathcal{I}_B)$. Grouping the two terms to the left, this is equal to $(\mathcal{P}_A \rightarrow \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \otimes \mathcal{P}_B)$. From Eq. (33), this further simplifies into $\overline{\mathcal{P}}_A \otimes \mathcal{P}_B$. $\square$

Going back to the physical interpretation of (56a), the pair of conditions (58) each corresponds to a term in the right-hand side of the equation. As these are linked by a ‘cap’, this reveals the set $\overline{\mathcal{A}} \otimes \mathcal{B}$ as the set of transformations that are compatible with no signaling from system A to system B **and** from B to A at the

same time; it is no signaling in both directions, whence the name. The same way, the set $\mathcal{A} \rightarrow \mathcal{B}$ is the one of transformations that respect no signaling from A to B **or** B to A. Thus, it may allow linear combinations of terms allowing signaling in each direction and consequently to indefinite causal order. The sets $\overline{\mathcal{A}} \prec \mathcal{B}$ and $\overline{\mathcal{A}} \succ \mathcal{B}$ lie in between, as they permit signaling in only one direction. Their elements indeed obey condition (57) but only one of the conditions (58). This discussion underlies the following chain of inclusions:

$$\overline{\mathcal{P}}_A \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}}_A \prec \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B, \quad (60a)$$

$$\overline{\mathcal{P}}_A \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}}_A \succ \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B. \quad (60b)$$

See Fig. 7 for a diagrammatic interpretation.

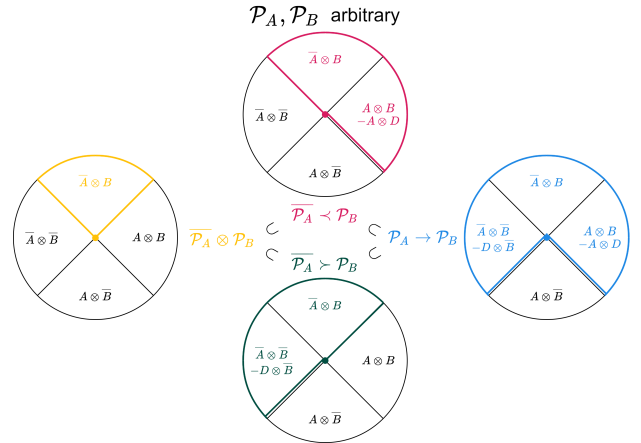


Figure 7: Inclusions (60) between the supports of the four bipartite composite projectors. These are the compositions involved in the characterization of the three kinds of transformations, Props. 2, 4, and 5 (in, respectively, blue, pink, and yellow). The blue part corresponds to the span of operators obeying condition (57), the pink (resp., green) part to the one obeying (57) and (58b) (resp., (58a)), and the yellow part to the one obeying (57), (58a) and (58b). See also the discussion around Fig. 9 in App. A.5.2.

Remark: defining the tensor product as no signaling. It can be noticed that the characterization of the no signaling transformation did not require M to be a valid transformation, i.e. an element of $\mathcal{P}_A \rightarrow \mathcal{P}_B$, only that it is a state structure. The no signaling composition of two state structures can be characterized as a corollary of Lemma 5 solely by requiring the two one-way signaling conditions to be valid simultaneously.

Corollary 2. *The set of trace-normalised positive operators $W \in \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ obeying the following two no signaling conditions: $\forall N_A \in \overline{\mathcal{A}}, \forall N_B \in \overline{\mathcal{B}}$,*

$$\text{Tr}_A[(N_A \otimes \mathbb{1}) \cdot W] = \text{Tr}_A\left[\left(\frac{\mathbb{1}}{c_A} \otimes \mathbb{1}\right) \cdot W\right], \quad (61a)$$

$$\text{Tr}_B[(\mathbb{1} \otimes N_B) \cdot W] = \text{Tr}_B\left[\left(\mathbb{1} \otimes \frac{\mathbb{1}}{c_B}\right) \cdot W\right], \quad (61b)$$

is the state structure of Definition 5 characterized by Eqs. (20).

Proof. A set of trace-normalized positive operators belonging to the space $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ is characterized by the projector $\mathcal{I}_A \otimes \mathcal{I}_B$. Making it obey conditions (61) restricts its projector to $(\mathcal{I}_A \otimes \mathcal{I}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{I}_B)$ by Lemma 5. Computing the intersections yields the projector $\mathcal{P}_A \otimes \mathcal{P}_B$ and the definition is reached. $\square$

This observation has two implications: on the one hand, it justifies its nickname of ‘no signaling composition’. On the other hand, it provides a ‘transformation-free’ definition of the tensor product based only on the validity of Lemma 5. This is an important result for the characterization since it avoids cyclic reasoning: the characterization of the transformation in Prop. 2 requires the characterization of the tensor from Def. 5. But without Corollary 2, this definition is either assumed *by fiat* or else it comes from Prop. 5 whose proof requires Prop. 2.

5.4 The full algebra of projectors

With respect to the discussion in Section 4.2.2, the further addition of the prec connector to the algebra of projectors increases the number of projectors in the Boolean lattice $\mathcal{D} \subseteq \mathcal{P} \subseteq \mathcal{I}$ under operations $\cap, \cup$. Any $\mathcal{P}$ can now be decomposed

into base projectors $\mathcal{P}_A, \mathcal{P}_B, \dots$ combined using any sequence of operations $\{\bar{\cdot}, \otimes, \prec, \rightarrow\}$ instead of $\{\bar{\cdot}, \otimes, \rightarrow\}$. As is the case for all operations on projectors defined so far (see Lemmas 1 and 2), composing using the prec yields a projector onto an operator system that preserves commutation; see App. A.5 for the proof. Summarizing the discussion of App. A, not only the six operations $\{\bar{\cdot}, \cap, \cup, \otimes, \prec, \rightarrow\}$ studied in this paper each preserve commuting projectors on operator systems, but any combination of these operations also does, leading to the following:

Theorem 1. *Let $\{\mathcal{P}_A, \mathcal{P}'_A, \mathcal{P}''_A, \dots\}$ be a set of commuting projectors on operator systems on a Hilbert space of operators $\mathcal{L}(\mathcal{H}^A)$. Let $\{\mathcal{P}_B, \mathcal{P}'_B, \mathcal{P}''_B, \dots\} \in \mathcal{L}(\mathcal{H}^B)$ be a similarly defined set. Let $\mathcal{P}$ be an expression built using a combination of elements from these two sets under the rules $\{\bar{\cdot}, \cap, \cup, \otimes, \prec, \rightarrow\}$ so that $\mathcal{P}$ is a projector on $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$. Let $\mathcal{P}'$ be another such expression. Then $\mathcal{P}$ and $\mathcal{P}'$ are a pair of projectors on operator systems of $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ that commute.*

The proof is straightforward. Its details are contained in App. A where each individual operation in $\{\bar{\cdot}, \cap, \cup, \otimes, \prec, \rightarrow\}$ is shown to map commuting projectors on operator systems to commuting projectors on operator systems, and so do their combinations.

When the base projectors associated with each tensor factor are fixed, the composite projectors are classified by their inclusion relations using $\cap, \cup$. As explained in Sec. 4.2.2, they form a bounded lattice, meaning that there is a finite number of possible sets of higher-order objects for a fixed number of parties, and that each set is either contained in or contains other sets, with the largest set being the fully signaling superset and the smallest being the no signaling subset. With respect to that, Eqs. (56) show that these lattices (as in Definition 7) can be extended to contain new, finer-grained elements built using the prec. In other words, Lemma 4 can be generalized to include elements built using the prec connector, without changing its top or bottom element. This extended lattice structure of projectors is summarised in the following:

Proposition 6 (Properties of the lattice of projectors). *Let $\{\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K\}$ be a set of k pro-*

jectors on operator systems, each associated with a specific Hilbert space $\mathcal{H}^A, \mathcal{H}^B, \dots, \mathcal{H}^K$, i.e. a set of base projectors.

General case: Any expression Γ built using each element of this set once and under the set of rules $\{\bar{\cdot}, \cap, \cup, \otimes, \prec, \rightarrow\}$ is a projector on an operator system over $\mathcal{L}(\mathcal{H}^A \otimes \dots \otimes \mathcal{H}^K)$. Moreover, it belongs to a lattice closed under operations $\{\cap, \cup\}$ such that

$$\mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K \subseteq \Gamma \subseteq \mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K. \quad (62)$$

Restricted case: Any expression Γ' built using each element of this set once and under the restricted set of rules $\{\otimes, \prec, \bar{\cdot} \rightarrow \cdot\}$ belongs to a sub-lattice such that

$$\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K \subseteq \Gamma' \subseteq \overline{\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K}. \quad (63)$$

The general case follows as a corollary of Theorem 1, whereas the restricted one is one of Lemma 3 under Eq. (60). A proof is provided in App. D.6. Notice that Eqs. (60) characterizes the simple example of a no signaling sub-lattice (as in the restricted case) built using the base projectors $\{\bar{\mathcal{P}}_A, \mathcal{P}_B\}$.

The fact there can be more than one base projector associated with each subsystem is what makes the algebra versatile. Each subsystem can indeed be itself composite, for example $\mathcal{L}(\mathcal{H}^A) = \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$; in that case, the various commuting projectors $\{\mathcal{P}_A, \mathcal{P}'_A, \mathcal{P}''_A, \dots\}$ to consider as bases can be for example the elementary one, $\mathcal{P}_A = \mathcal{P}_{A_0} \otimes \mathcal{P}_{A_1}$, the one-way signaling from A_0 to A_1 one, $\mathcal{P}'_A = \mathcal{P}_{A_0} \prec \mathcal{P}_{A_1}$, the two-way signaling one, $\mathcal{P}''_A = \bar{\mathcal{P}}_{A_0} \rightarrow \mathcal{P}_{A_1}$, etc.

As mentioned in Sec. 4.1, nothing forbids *a priori* to consider *any* family of commuting projectors on operator system as possible choices of bases for a subsystem –not only the ones built from $\mathcal{I}$ and the algebra rules $\{\bar{\cdot}, \otimes, \prec, \rightarrow\}$. The biased quantum theory example of Sec. F is an example of such a choice of arbitrary projectors as base projectors. In this paper, and especially in the applications of the next section, we focus on families built from $\mathcal{I}$ using the rules of the algebra only because of their interpretation as (structure-preserving) quantum supermaps with constraints on their signaling structure. Yet, other physically-relevant choices can be made, for example building from the family of bases $\{\mathcal{I}, \Delta\}$,

where Δ projects onto diagonal operators w.r.t. a fixed basis. In that case, $\mathcal{I}$ defines the states of quantum theory and Δ its ‘dephased’ subspace, i.e. classical states. This direction is left open for future research (a longer discussion is provided in App. A.1).

Remark: the algebra of projectors is (almost) BV/Rétoré logic. We pointed out the connection with MALL in the remark in Sec. 4.2.2 and in App. A.4.3. With the further addition of the prec, the algebra defined by the above theorem can be seen as an instance of BV- or Rétoré logic, as first noticed in Ref. [26] (although their choice of additives connectives is different from ours).

5.5 The algebra of projectors recovers no signaling

In this section, we observed that the notion of no signaling could be generalized for the case of state structures. In particular, it was noticed that the notion could be encoded by a new rule $\prec$ in the algebra of projectors, introduced under the name prec.

Bipartite state structures obtained by combining $\mathcal{A}$ and $\mathcal{B}$ fall into four classes, essentially determined by conditions (57), (58b), and (58a) represented graphically in, respectively, Figs. 6a, 6b, and 6f. All of these classes must satisfy the first condition, which amounts to being a valid transformation and to which corresponds the projector $\mathcal{P}_A \rightarrow \mathcal{P}_B$. If in addition condition (58b) (respectively, (58a)) is satisfied, the set is restricted to a transformation that is no signaling from the output to the input (resp., from the input to the output), to which corresponds the projector $\bar{\mathcal{P}}_A \prec \mathcal{P}_B$ (resp., $\bar{\mathcal{P}}_A \succ \mathcal{P}_B$). If both conditions are satisfied, the set is no signaling, to which corresponds the projector $\bar{\mathcal{P}}_A \otimes \mathcal{P}_B$.

Therefore, working with projectors provides a handy way of assessing the signaling structure within a state structure.

6 Implications of characterizing signaling

With the formalization of the projective characterization and the introduction of the one-way signaling subsets using the prec, we now have the

tools to address signaling in higher-order quantum transformations.

6.1 Algebraic manipulations using the prec: two birds with one stone

On purely algebraic grounds, the transformation $\rightarrow$ has two defining properties that make it inconvenient to work with. The first is not commuting with negation¹²:

$$\overline{\mathcal{P}_A \rightarrow \mathcal{P}_B} \neq \overline{\mathcal{P}_A} \rightarrow \overline{\mathcal{P}_B}. \quad (64)$$

This leads to difficulties in relating seemingly similar sets together using the algebra alone. Two expressions involving the same states and effects may be incomparable because one features a negation on several subsystems at once while the other does not. For example, $\overline{\mathcal{P}_A \rightarrow \mathcal{P}_B} = \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}$ cannot be compared with $\overline{\mathcal{P}_A} \rightarrow \overline{\mathcal{P}_B} = \overline{\mathcal{P}_A} \otimes \mathcal{P}_B$ without explicitly doing a computation like (131) to show set inclusion.

The second inconvenient property is that $\rightarrow$ is not associative:

$$\mathcal{P}_A \rightarrow (\mathcal{P}_B \rightarrow \mathcal{P}_C) \neq (\mathcal{P}_A \rightarrow \mathcal{P}_B) \rightarrow \mathcal{P}_C. \quad (65)$$

While the algebraic difficulties of non-associativity can be alleviated by replacing the transformation connector “ $\rightarrow$ ” with the associative “ $\overline{\cdot} \rightarrow \cdot$ ”, this issue still impedes the interpretation of type formulae. A naive interpretation of the arrow as a transformation may lead an unsuspecting reader to miss comparable state structures. For example, consider types $((A_0 \rightarrow A_1) \rightarrow A_2) \rightarrow A_3$ and $(A_1 \rightarrow A_2) \rightarrow (A_0 \rightarrow A_3)$, the first formula depicts two cascaded transformations whose input is a channel from A_0 to A_1 , whereas the second is a single transformation whose input is a channel from A_1 to A_2 . It is tempting to conclude they have nothing in common. However, some computations will reveal that they both feature $\overline{\mathcal{P}_{A_0}} \otimes \mathcal{P}_{A_1} \otimes \overline{\mathcal{P}_{A_2}} \otimes \mathcal{P}_{A_3}$ as their no

signaling subset (as in Def. 7), and accordingly they should be comparable. If the base types in this example are the quantum states, then the two expressions are more than comparable: they describe the same set of *quantum 2-combs* [3]! Two expressions involving the same no signaling subset may be incomparable because there is no way to ‘shuffle the parentheses’ to put their constituents in the same order.

Surprisingly, the prec does not lead to these issues, as it commutes with the negation:

$$\overline{\mathcal{P}_A \prec \mathcal{P}_B} = \overline{\mathcal{P}_A} \prec \overline{\mathcal{P}_B}. \quad (66)$$

And, like the tensor, it is associative:

$$\begin{aligned} (\mathcal{P}_A \prec \mathcal{P}_B) \prec \mathcal{P}_C &= \mathcal{P}_A \prec (\mathcal{P}_B \prec \mathcal{P}_C) \\ &= \mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C. \end{aligned} \quad (67)$$

(See Eqs. (162) and (163) in App. A.5.1 for the proofs.) With these two obstacles gone, it becomes possible to compare state structures simply by algebraic manipulations of their projectors. As the tensor and the transformation split into precs, Eqs. (56), and since the algebra obeys the De Morgan relations (42), the ‘impractical’ tensor and transformation connectors can be expressed as precs instead. Any projector $\mathcal{P}$ built as in Theorem 1 and whose decomposition into base projectors features $\otimes$ and $\rightarrow$ can be rewritten into unions and intersections of expressions built using the prec alone.

Now, in the algebra, the cap distributes over the cup,

$$(\mathcal{P}_A \cup \mathcal{P}'_A) \cap \mathcal{P}''_A = (\mathcal{P}_A \cap \mathcal{P}''_A) \cup (\mathcal{P}'_A \cap \mathcal{P}''_A). \quad (68)$$

(See Eq. (105) in App. A.2.) Furthermore, the cap and the cup both distribute over the prec,

$$(\mathcal{P}_A \cup \mathcal{P}'_A) \prec \mathcal{P}_B = (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}'_A \prec \mathcal{P}_B), \quad (69a)$$

$$\mathcal{P}_A \prec (\mathcal{P}_B \cap \mathcal{P}'_B) = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \prec \mathcal{P}'_B), \quad (69b)$$

where the above are also satisfied when caps and cups are switched, $\cap \leftrightarrow \cup$. (See App. A.5.3.) Consequently, distributivity can be used to define a normal form: expressions involving these three connectors can always be rewritten so that they are put into a union of intersections of one-way signaling compositions. As these three operations

¹²In categorical terms, this is a manifestation of the $\ast$ -autonomous character of the category of higher-order transformations [9]. If the $\rightarrow$ were to commute with the negation, the category would be compact closed instead of $\ast$ -autonomous. This is what happens in the case of quantum theory where the projectors on both sides are the identity: $\overline{\mathcal{I}_A} \otimes \overline{\mathcal{I}_B} = \overline{\mathcal{I}_A} \otimes \overline{\mathcal{I}_B} = \mathcal{D}_A \otimes \mathcal{D}_B$ and which explains why the bipartite deterministic effects of quantum theory are not two-way signaling. See the discussion in Sec. 6.2 and in App. B.

are associative, the intermediate parentheses can be dropped everywhere.

For example, in the bipartite case, there are eight possible one-way signaling compositions, or ‘prec chains’: $\mathcal{P}_A \prec \mathcal{P}_B$, $\mathcal{P}_A \succ \mathcal{P}_B$, $\overline{\mathcal{P}}_A \prec \mathcal{P}_B$, etc. The claim is that expressions built using $\{\neg, \cap, \cup, \prec\}$ can be rewritten using the distribution rules into unions of intersections of these eight prec chains like e.g., $\mathcal{P}_A \succ \overline{\mathcal{P}}_B$ or $(\mathcal{P}_A \prec \mathcal{P}_B) \cup ((\overline{\mathcal{P}}_A \prec \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \succ \overline{\mathcal{P}}_B))$, and that expressions not in this form, like e.g., $((\mathcal{P}_A \prec \mathcal{P}_B) \cup (\overline{\mathcal{P}}_A \prec \mathcal{P}_B)) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B)$, can be ‘simplified’ through distribution until it reaches a ‘union of intersections of prec chains’, e.g., $((\mathcal{P}_A \prec \mathcal{P}_B) \cup (\overline{\mathcal{P}}_A \prec \mathcal{P}_B)) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B)$ is rewritten into $((\mathcal{P}_A \prec \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B)) \cup ((\overline{\mathcal{P}}_A \prec \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B))$.

Definition 9 (Normal form). *Let $\mathcal{P}$ be a projector on an operator system built from a set of k base projectors $\{\mathcal{P}_A, \mathcal{P}_B, \dots, \mathcal{P}_K\}$ composed together under operations $\{\neg, \cap, \cup, \otimes, \prec, \succ\}$ so that $\mathcal{P}$ acts on $\mathcal{L}(\mathcal{H}^A \otimes \dots \otimes \mathcal{H}^K)$. Then, a **normal form** Γ of $\mathcal{P}$ is a projector equivalent to $\mathcal{P}$ obtained as unions of intersections of expressions built from operations $\{\neg, \prec\}$ alone:*

$$\Gamma = \bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \tilde{\mathcal{P}}_{\sigma_{ij}(B)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)} \right), \quad (70)$$

in which there are x unions of expressions labeled by index i , and each expression involves y_i intersections of sub-expressions labelled by index j , where σ_{ij} is an element of the permutation group on k element, like e.g. $\sigma_{00}(A) = A, \sigma_{01}(A) = B, \dots$, so that each $\tilde{\mathcal{P}}_{\sigma_{ij}(A)}$ is a choice of a base projector which is potentially negated depending on indices i and j . Each sub-expression is thus a permutation of $\tilde{\mathcal{P}}_A \prec \tilde{\mathcal{P}}_B \prec \dots \prec \tilde{\mathcal{P}}_K$ where the position of the negations depends on i and j (note that the indices do not necessarily run over the full permutation group).

Remark that the convention of ordering expressions in alphabetical label order is dropped in the normal form. It is replaced by an ordering that allows reading prec chains from left to right using associativity. In other words, normal forms only use the $\prec$; its reversed version, $\succ$, must also be dropped. For example, the normal form of $(\mathcal{P}_A \prec \mathcal{P}_B) \succ \mathcal{P}_C$ is written as $\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B$.

Theorem 2. *Any projector $\mathcal{P}$ as in the general case of Prop. 6 has a normal form.*

Proof. By Proposition 6, it is a projector on operator system. By Eqs. (56), the projector can be put in a form involving operations $\{\neg, \cap, \cup, \prec\}$ only. The normal form can always be reached by first distributing negations over intersections and unions using De Morgan law (30) and over the prec using Eq. (66), then by distributing the prec over unions and intersections using (69), and finally by distributing the caps over the cups using (68). This procedure is explicitly proven in App. D.7. $\square$

This implies that any expression can be brought down to the union of several one-way signaling expressions (which further decompose into intersections if needed). This way of rewriting exists for *any* state structure, lifting any ambiguity induced by the negation or non-associativity. In the first example discussed above, the expressions reduce to $\overline{\mathcal{P}}_A \rightarrow \mathcal{P}_B = (\mathcal{P}_A \prec \overline{\mathcal{P}}_B) \cap (\mathcal{P}_A \succ \overline{\mathcal{P}}_B)$ and $\overline{\mathcal{P}}_A \rightarrow \overline{\mathcal{P}}_B = \overline{\mathcal{P}}_A \rightarrow \overline{\mathcal{P}}_B = (\mathcal{P}_A \prec \overline{\mathcal{P}}_B) \cup (\mathcal{P}_A \succ \overline{\mathcal{P}}_B)$. Now one can meaningfully compare them: they are both built from a combination of the same two one-way signaling transformations, but one set is the intersection and the other is the union, we thus conclude that $\overline{\mathcal{A}} \rightarrow \mathcal{B} \subseteq \overline{\mathcal{A}} \rightarrow \overline{\mathcal{B}}$ simply by inspecting the normal form of their respective projectors.

Moreover, the prec connector and the decomposition it induces also allow for an identification of the possible signaling directions these two sets may feature. Indeed, since the prec is exactly the one-way signaling condition, the above normal form is a breakdown of a general expression into several fully ordered expressions, or ‘prec chains’. That is, expressions built using only the prec like $\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C$ characterizes a set of objects with a fixed signaling order. Indeed, since $\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C = (\mathcal{P}_A \prec \mathcal{P}_B) \prec \mathcal{P}_C$, C cannot signal to A and to B , and B cannot signal to A ; in the same time since $\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C = \mathcal{P}_A \prec (\mathcal{P}_B \prec \mathcal{P}_C)$, A might signal to B and C while B might signal to C ¹³. These chains

¹³More precisely: certain operators W_{ABC} in $\mathcal{A} \prec \mathcal{B} \prec \mathcal{C}$ may allow A to signal *deterministically* to B and to C by suitably choosing the (local deterministic) effect $N_A \in \overline{\mathcal{A}}$ she will apply on W_{ABC} . And the same way, B may signal to C by suitably choosing his effect $N_B \in \overline{\mathcal{B}}$, but he will

are then combined first by requiring no signaling between some subsystems, i.e. using the cap, and then by allowing signaling in both directions, i.e. using the cup. For example, to further impose no signaling between A and B in $\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C$, one intersects it with $(\mathcal{P}_A \succ \mathcal{P}_B) \prec \mathcal{P}_C$ so that $(\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C) \cap (\mathcal{P}_B \prec \mathcal{P}_A \prec \mathcal{P}_C) = (\mathcal{P}_A \otimes \mathcal{P}_B) \prec \mathcal{P}_C$. Then, to loosen the requirement that C can signal to neither A nor B , one takes the union of it with $(\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_C \prec \mathcal{P}_B \prec \mathcal{P}_A) = (\mathcal{P}_A \otimes \mathcal{P}_B) \succ \mathcal{P}_C$. The obtained projector, $\overline{\mathcal{P}_A \otimes \mathcal{P}_B} \rightarrow \mathcal{P}_C$, has the normal form $((\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C) \cap (\mathcal{P}_B \prec \mathcal{P}_A \prec \mathcal{P}_C)) \cup ((\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_C \prec \mathcal{P}_B \prec \mathcal{P}_A))$. This indeed defines a set of objects in which A and B cannot signal to each other, but both might signal to and receive signaling from some party C , and all this information can be read directly from its normal form.

Nevertheless, because the normal form encodes the possible signaling directions, it is not unique. Indeed, as the prec is associative, this means that the impossibility to signal is transitive, so that in an expression like $\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C$, C cannot signal to A . This implies that some formulae are redundant. For example, the following can be shown in the algebra:

$$\begin{aligned} (\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C) \cap (\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B) = \\ (\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C) \cap (\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B) \\ \cap (\mathcal{P}_A \prec \mathcal{P}_C \prec \mathcal{P}_B). \end{aligned} \quad (71)$$

This is intuitively straightforward: if C cannot signal to A and B and vice-versa, this also includes the situation in which C cannot signal to A while B cannot signal to C . Therefore, the third term on the right-hand side of the above is redundant, and these two normal forms are equivalent. Improving the definition of the normal form so as to incorporate this redundancy would hopefully make the normal form unique. This is left open as a future research direction.

Hence, putting a projector in normal form provides a direct way to read the possible signaling structures of the set of higher-order transformations it represents. Moreover, putting every projector under consideration into normal form provides a quick way to spot equivalent state structures. Nevertheless, having equivalent normal forms is but a sufficient condition for the never be able to signal to A , no matter his choice of N_B .

equivalence of projectors. One should complete the test by actually computing the intersection (or union) of projectors with inequivalent normal forms to reach a definitive conclusion. For example, in Eq. (71) the normal form allows one to regroup the terms into two projectors $\mathcal{P}_1 := (\mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C) \cap (\mathcal{P}_C \prec \mathcal{P}_A \prec \mathcal{P}_B)$ and $\mathcal{P}_2 := \mathcal{P}_A \prec \mathcal{P}_C \prec \mathcal{P}_B$. From there, it is obvious that the right-hand side is contained in the left-hand side since this equation has the form $\mathcal{P}_1 \subseteq \mathcal{P}_1 \cap \mathcal{P}_2$ and it remains to show that $\mathcal{P}_1 = \mathcal{P}_1 \cap \mathcal{P}_2$ by explicit computation, which is also made simpler since one knows that they do not have to explicitly compute the intersection contained in $\mathcal{P}_1$ as it appears on both side of the equation.

A concrete example of the use of the normal form, as well as of an equivalence of normal form, is provided in App. G.

Summarizing, the algebraic fact that the prec is associative and commutes with the negation allows one to write a ‘normal’ form useful to compare types of transformations. This gives a tool to determine whether two types of higher-order transformations are equivalent by sole inspection of their projectors. In addition, the physical fact that the prec is a one-way signaling composition, combined with the algebraic rules, allows one to split projectors into several terms with a fixed signaling direction between their base state structures. This implies that one can know the possible signaling structure that a higher-order theory may feature by sole inspection of its projector.

6.2 When quantum combs are isomorphic to quantum networks

The next part of the results is concerned with the relationship between two different ways of building higher-order objects using the developed formalism of projectors. This will be shown to be an example for which the normal form allows to swiftly recover results of Ref. [5].

Following Ref. [3], a **network** is defined as the causally ordered (i.e. one-way signaling) succession of ‘nodes’ of the same state structure. A ‘1-network of base $\mathcal{A}$ ’ will be the set $\mathcal{A} \subset \mathcal{L}(\mathcal{H}^A)$ itself, thereafter noted with an index as $\mathcal{A}_0$ to distinguish between the multiple copies of the state structure $\mathcal{A}$ instantiated on an increasingly larger number of systems. The ‘2-network of base $\mathcal{A}$

will be the set $\mathcal{A}_0 \prec \mathcal{A}_1 \subset \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$, the ‘3-network’ will be $\mathcal{A}_0 \prec \mathcal{A}_1 \prec \mathcal{A}_2$, etc. up to the ‘ n -network’ defined as $\mathcal{A}_0 \prec \mathcal{A}_1 \prec \dots \prec \mathcal{A}_{n-1} \subset \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1} \otimes \dots \mathcal{H}^{A_{n-1}})$. A common occurrence of this structure is the network whose base is a quantum channel so that $\mathcal{A}$ is characterized by a projector $\mathcal{P}_{\mathcal{A}} = \mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}$; it is called a **quantum network** (notice that it will require twice as many systems since the base is defined as a state structure involving two subsystems). This quantum network, here associated with some party that will be called Alice, represents the successive operations of that party. If Alice has a single-node quantum network, it means that Alice acts once on subsystem A_0 with a quantum channel and outputs a subsystem A_1 in a quantum state defined in space $\mathcal{L}(\mathcal{H}^{A_1})$. If she has a network with two nodes, she will act a first time on A_0 and output a first state at A_1 , then a second time on A_2 , now potentially using any size of ancillary qudit as a memory register she preserved from her first operation, and outputs a second state in A_3 . And so on for all numbers of nodes, as defined recursively.

Another way of building a higher-order state structure is the **comb**, which consists of recursively transforming into a base type: the ‘1-comb of base $\mathcal{A}$ ’ is again $\mathcal{A}_0$, then the 2-comb is $\mathcal{A}_0 \rightarrow \mathcal{A}_1$, the 3-comb is $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \rightarrow \mathcal{A}_2$, etc up to the n -comb defined as $(\dots(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \rightarrow \dots) \rightarrow \mathcal{A}_n \subset \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1} \otimes \dots \mathcal{H}^{A_{n-1}})$. As is the case for the network case, a common occurrence of this structure is the comb whose base is a quantum channel, called the **quantum comb**.

When the two constructions were introduced in Ref. [3], it was proven that a quantum network *is* a quantum comb. When treated using the formalism developed here, there is a stark contradiction. All quantum combs are built using the transformation, $\rightarrow$, which permits two-way signaling. How could it be that they are all equivalent to networks which are built using the prec – that is, objects featuring a single direction of signaling? Besides, why is the 1-comb (built using the two-way signaling transformation) equivalent to the quantum channel (which is causal)? Why does the quantum 2-comb reduce to a two-node quantum network, i.e. a map acting on two quantum states, when by definition it should be a supermap, i.e. a map acting on a quantum channel? And why does it reduce to a succession of two

operations that have a well-defined direction of signaling between parties, and not, as its projector suggests, a nesting of two two-way signaling compositions, resulting in four possible signaling directions?

To phrase this issue formally, some notation is required. Let $\mathcal{A}_i^{\text{state}} \subset \mathcal{L}(\mathcal{H}^{A_i})$ be the state structure of states on system A_i so that its projector is $\mathcal{I}_{A_i}$, and let $\mathcal{A}_{i,j}^{\text{channel}} \subset \mathcal{L}(\mathcal{H}^{A_i} \otimes \mathcal{H}^{A_j})$ be the state structure of channels between systems A_i and A_j so that its projector is $\mathcal{I}_{A_i} \rightarrow \mathcal{I}_{A_j}$. Writing $\mathcal{A}_{\text{base}}$ as either of these two structures used as the base of a larger state structure (a name appearing in the index indicates that the labeling of the subsystems has not been specified yet), the network and comb can be abstractly defined using recursive formulae. Let $\mathcal{A}_{1\text{-network}} = \mathcal{A}_{1\text{-comb}} = \mathcal{A}_{\text{base}}$. Then, the definitions of Ref. [3] can be abstracted as $\mathcal{A}_{n\text{-network}} \equiv \mathcal{A}_{(n-1)\text{-network}} \prec \mathcal{A}_{\text{base}}$ and $\mathcal{A}_{n\text{-comb}} \equiv \mathcal{A}_{(n-1)\text{-comb}} \rightarrow \mathcal{A}_{1\text{-comb}}$. Finally, let the labeling of the base state structures be chosen such that the subsystems in each state structure coincide.

The structures we want to compare are the quantum networks, the quantum combs, and the combs of quantum states with twice as many nodes. In symbols, the state structure of quantum networks (i.e. of networks whose base is quantum channels) with n nodes reads $\mathcal{A}_{0,1}^{\text{channel}} \prec \dots \prec \dots \prec \mathcal{A}_{2n-2,2n-1}^{\text{channel}}$ which, in the language of Refs. [4, 5], carries the type $(A_0 \rightarrow A_1) \prec \dots \prec (A_{2n-2} \rightarrow A_{2n-1})$ (where each A ’s are of the type of quantum states), and it is associated with the projector

$$\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-network})} \equiv (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \prec \dots \prec (\mathcal{I}_{A_{2n-2}} \rightarrow \mathcal{I}_{A_{2n-1}}). \quad (72)$$

The state structure of quantum combs (i.e. of combs whose base are quantum channels) with n nodes reads $(\dots(\mathcal{A}_{n-1,n}^{\text{channel}} \rightarrow \mathcal{A}_{n-2,n+1}^{\text{channel}}) \dots) \rightarrow \mathcal{A}_{0,2n-1}^{\text{channel}}$, it carries the type $(\dots((A_{n-1} \rightarrow A_n) \rightarrow (A_{n-2} \rightarrow A_{n+1})) \rightarrow \dots) \rightarrow (A_0 \rightarrow A_{2n-1})$, and it is associated with the projector

$$\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})} \equiv (\dots(\mathcal{I}_{A_{n-1}} \rightarrow \mathcal{I}_{A_n}) \rightarrow \dots) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_{2n-1}}). \quad (73)$$

The state structure of combs based on states with $2n$ nodes reads

$(\dots((\mathcal{A}_0^{\text{state}} \rightarrow \mathcal{A}_1^{\text{state}}) \rightarrow \mathcal{A}_2^{\text{state}}) \dots) \rightarrow \mathcal{A}_{2n-1}^{\text{state}}$; it carries the type $(\dots((A_0 \rightarrow A_1) \rightarrow A_2) \rightarrow \dots) \rightarrow A_{2n-1}$; it is associated with the projector

$$\mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})} \equiv (\dots(\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \rightarrow \dots) \rightarrow \mathcal{I}_{A_{2n-1}}. \quad (74)$$

To answer the above interrogations, it can be found in the literature 1) that quantum channels are no signaling from output to input, and are equivalent to quantum 1-network and quantum 1-combs; 2) that the first two state structures defined in (72) and (73) above are equivalent [3, Theo. 8]; 3) that the elements M of these state structures can be characterized in the CJ picture by the causality condition [3, Theo. 5]:

$$\forall \in 1, \dots, n,$$

$$\begin{aligned} \text{Tr}_{A_{2i+1}} [M^{(i)}] &= \frac{1}{d_{A_{2i}}} \text{Tr}_{A_{2i} A_{2i+1}} [M^{(i)}] \otimes \mathbb{1}_{A_{2i}}, \\ M^{(n)} &\equiv M, \quad M^{(j)} \equiv \text{Tr}_{A_{2j} \dots A_{2n-1}} [M], \quad j < n; \end{aligned} \quad (75)$$

4) that the last two state structures defined in (73) and (74) above are equivalent [5, Prop. 6].

We now recover and explain these results using the projective characterization. These four statements can indeed be proven simply by algebraic manipulations of the projectors. They amount to proving that 1) $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} = \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}$; 2) Equations (72) and (73) are equivalent, i.e. $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})} = \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-network})}$; 3) Equation (75) is an equivalent way of defining the validity subspace of quantum networks, Eq. (72); 4) Equations (73) and (74) are equivalent, i.e. $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})} = \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})}$.

Transformations between quantum states.

The case $n = 1$ is treated as a warm-up before the general case. Equations (73), (72) and (74) all reduce to $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}$, so items 2) and 4) hold. What remains to be proven are items 1) and 3) in the single-node case. The following lemma does that, and it underlies the proof for the general case.

Lemma 6. *The transformation operation on projectors, $\rightarrow$, simplifies into a prec operation when either of the projectors is the identity. The prec operation on projectors, $\prec$, simplifies into a tensor operation when either of the projectors is*

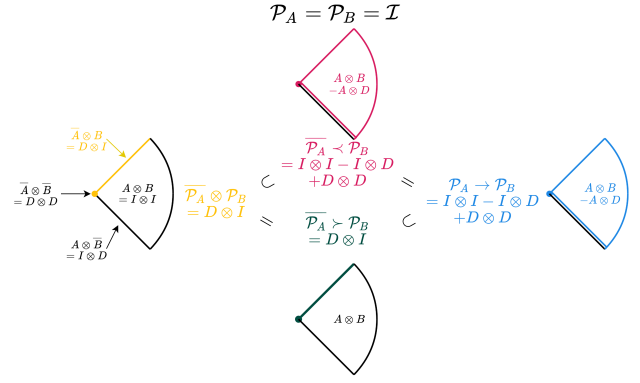


Figure 8: Diagrammatic representation of Lemma 6: isomorphisms between the compositions when the base projectors are the identity, $\mathcal{I}$, i.e. in the case of transformations between quantum states. Compared to the general situation, Fig. 7, the A-to-B one-way signaling composition becomes equivalent to the two-way one, whereas the B-to-A one-way signaling composition becomes equivalent to the no signaling one.

depolarizing. In equation,

$$\begin{aligned} \mathcal{P}_A \rightarrow \mathcal{P}_B &= \bar{\mathcal{P}}_A \prec \mathcal{P}_B \\ \iff \mathcal{P}_A &= \mathcal{I}_A \text{ or } \mathcal{P}_B = \mathcal{I}_B; \end{aligned} \quad (76a)$$

$$\begin{aligned} \bar{\mathcal{P}}_A \prec \mathcal{P}_B &= \bar{\mathcal{P}}_A \otimes \mathcal{P}_B \\ \iff \mathcal{P}_A &= \mathcal{D}_A \text{ or } \mathcal{P}_B = \mathcal{D}_B. \end{aligned} \quad (76b)$$

(This follows from the definitions of the operations, see App. B for the proof.)

The first equation of this lemma implies that transformations between quantum states, characterized by $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}$, are equivalent to a causal succession of a functional and a state, $\mathcal{D}_{A_0} \prec \mathcal{I}_{A_1}$. This proves statement 1), $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} = \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}$. Remark these actually are two different ways of defining the quantum channel: as the most general transformation that takes a quantum state to a quantum state, or as a measurement coherently followed by a preparation of a quantum state¹⁴ (this is also mentioned in example E.3).

As for the second equation of the lemma, it is for example the reason why the single-partite process matrix reduces to an effect and a state in

¹⁴By coherent, it is meant ‘as any element of the convex hull of measure-and-reprepare maps’, so this set does not represent only entanglement breaking maps, i.e. of the form $\mathcal{M}(\bullet) = \rho \text{Tr}[\sigma \cdot \bullet]$, but any linear map $\mathcal{M}(\bullet) = \sum_i \rho_i \text{Tr}[\sigma_i \cdot \bullet]$. Recall that in this non-CJ representation, the CPTP condition comes as constraints on the ρ_i ’s and σ_i ’s.

tensor product [8]: the 1-PM is a functional on channels, characterized by the negated projector $\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}}$. Using the first part of the lemma, $\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}} = \overline{\mathcal{D}_{A_0} \prec \mathcal{I}_{A_1}}$. The negation is distributed over a prec, $\overline{\mathcal{D}_{A_0} \prec \mathcal{I}_{A_1}} = \mathcal{I}_{A_0} \prec \mathcal{D}_{A_1}$, and finally the second part of the lemma is used, $\mathcal{I}_{A_0} \prec \mathcal{D}_{A_1} = \mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}$. These algebraic manipulations on projectors quickly led to the conclusion that the functionals on quantum channels, characterized by $\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}}$, are equivalent to a quantum state on the input system followed by a measurement applied on its output, $\mathcal{I}_{A_0} \prec \mathcal{D}_{A_1}$, and that the state and measurement are causally disconnected, $\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}$.

More generally, the first equation (76a) states that any transformation from (or to) an arbitrary state structure to (or from) the state structure of quantum states will automatically be no signaling from output to input. Contrastingly, the second equation (76b) states that any way of combining a state structure with the single element ‘quantum measurement’ state structure, i.e. $\{\mathbb{1}\}$, is automatically no signaling from and to it. These relations are depicted in Fig. 8 for the case where both A and B are quantum states. This should be compared with the general case of Fig. 7.

Finally, as previously discussed in the remark of Sec. 5.1, Eq. (76a) is the case where the definition of no signaling, Eq. (44), overlaps with the one of a transformation, see Eq. (46), so the usual definition of no signaling is recovered. The fact that these two equations coincide in this case is exactly statement 3) in the $n = 1$ case: it suffices to notice that $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} = \mathcal{D}_{A_0} \prec \mathcal{D}_{A_1}$ and that this projector applied on bipartite operators is exactly enforcing the condition (75) in the $n = 1$ case, i.e. $\forall M^{(1)} \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1}) : (\mathcal{D}_{A_0} \prec \mathcal{D}_{A_1})\{M^{(1)}\} = M^{(1)} \iff \text{Tr}_{A_1}[M^{(1)}] = \mathbb{1}_{A_0}$.

Equivalence of quantum combs and networks. We shall now generalize this equivalence for all n . Reinterpreting $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} = \mathcal{D}_{A_0} \prec \mathcal{I}_{A_1}$, the quantum channel can be seen as a sort of network in which the nodes alternate between an effect and a state (i.e., the wires alternate between pointing up and down). This was shown to imply (75) in the $n = 1$ case. Going to $n > 1$, a network of quantum channels is then an alternating network of effects and states as associativity can be used, and it indeed implies (75)

in general.

Theorem 3. *The projectors to a quantum network (72), to a quantum comb (73), and to a (twicely longer) comb based on quantum states (74) are all equivalent to the following projector, characterizing an alternating network of quantum effects and states:*

$$\mathcal{P}^{(2n)} = \overline{\mathcal{I}_{A_0}} \prec \mathcal{I}_{A_1} \prec \dots \prec \overline{\mathcal{I}_{A_{2n-2}}} \prec \mathcal{I}_{A_{2n-1}}. \quad (77)$$

Meaning that with suitable normalization, the state structure of networks of order n based on quantum channels is equivalent to the one of combs of order n based on quantum channels, and it is also equivalent to the one of combs of order $2n$ based on quantum states.

In addition, the elements M of the state structure defined by this projector obey equation (75).

The proof is presented in App. D.8. This theorem means that an operator M obeying the causality condition (75) can be interpreted as an element of four equivalent state structures: (1) from Eq. (73) as a valid quantum n -comb, i.e. (the CJ representation of) a map on $2n$ systems that transforms the recursively defined map on the $n - 1$ nodes split between the $2n - 2$ systems $A_1, \dots, A_{2n-2}$ into a single-node map, i.e. a quantum channel from A_0 to A_{2n-1} . (2) from Eq. (72) as a valid quantum network of channels with n nodes, i.e. (the CJ representation of) a map on $2n$ systems $A_0, \dots, A_{2n-1}$ that represents the n successive operations of a party on n nodes: at node j she applies a channel between systems A_{2j-2} and A_{2j-1} that can depend on (that is, use any size of ancillary memory from) the $j - 1$ previous nodes. (3) from Eq. (74) as a valid comb of states on $2n$ nodes, i.e. (the CJ representation of) a map on $2n$ systems that transforms the recursively defined map on the $2n - 1$ systems $A_0, \dots, A_{2n-2}$ into a quantum state on the $2n$ -th system A_{2n-1} . And (4) from Eq. (77) as a valid quantum network alternating between quantum measurements and states at each of the $2n$ nodes, i.e. (the CJ representation of) a map on $2n$ systems $A_0, \dots, A_{2n-1}$ that represents the $2n$ successive operations of a party on $2n$ nodes: at odd nodes, she measures a quantum state, at even nodes, she prepares one. Each operation is potentially conditioned by the previous nodes but is independent of the future ones.

Although these four definitions appear quite different, notice that equation (77) is actually the normal form of Eqs. (72), (73), and (74)! Hence the equivalence between seemingly unrelated definitions of higher-order maps is inferred simply by working out their normal form. The fact that the two comb-like definitions result in structures with a single signaling direction is now directly apparent, as the structure features a single ‘prec chain’ in normal form. As mentioned above, this is a very peculiar behavior: From equation (56b) one would have expected the normal forms to be the union of many prec chains with different signaling orderings. That is, the combs should typically have an indefinite causal order between their nodes. Lemma 6 auspiciously prevents that. The isomorphisms between $\prec$ and $\rightarrow$ in the case of composite state structure whose base structures are sets of quantum states explain this apparent counter-logical equivalence of one-way signaling objects (the networks of states) with various a priori two-way signaling ones. Quantum theory is thus very tame in the sense that successive and/or nested transformations of quantum states or channels happen to automatically have no signaling from output to input.

On the contrary, nesting any other state structures in a comb-like construction will in general result in indefinite causal order. An example of this in the case of Multi-round Process Matrices [11] is provided in App. G. It also features another example of the use of the normal form.

7 Discussion

In this work, a correspondence was developed between several characterizations of the theory of higher-order quantum transformations. Starting at the most abstract layer, its formulation in terms of types, we added a new layer, its formulation in terms of superoperator projectors, to bridge to a less abstract layer, its formulation in terms of sets of CJ operators – here called state structures. Correspondences between the characterizations can be summed up as follows: Given types A and B , types $\bar{A}$, $A \otimes B$, and $A \rightarrow B$ can be constructed using the semantic rules $\{1, (,), \rightarrow\}$. To these types correspond projectors: given projectors $\mathcal{P}_A$ and $\mathcal{P}_B$, projectors $\bar{\mathcal{P}}_A$, $\mathcal{P}_A \otimes \mathcal{P}_B$ and $\mathcal{P}_A \rightarrow \mathcal{P}_B$ can be constructed using the algebraic rules $\{\bar{\cdot}, \otimes\}$ as in equations (15c), (20c),

and (23c). To these types and projectors correspond state structures: given state structures $\mathcal{A}$ and $\mathcal{B}$ as in equation (3), the state structures $\bar{\mathcal{A}}$, $\bar{\mathcal{A}} \otimes \mathcal{B}$, and $\mathcal{A} \rightarrow \mathcal{B}$ can be constructed by requiring that, respectively, conditions (13), (61), and (57) hold.

The main novelty is the prec connector, Prop. 8, giving the state structure $\bar{\mathcal{A}} \prec \mathcal{B}$ of maps which are no signaling from the output to the input. This allows to speak about the possible signaling structure of the maps characterized by the projective framework. In addition to that, the properties of the prec hints towards a systematic way to compare different projective expressions as they allow to expand them as unions and intersections of expressions involving only this prec connector.

In the meantime, we identified the features of the algebra of projectors. We showed in particular that it was a Boolean algebra under the rules $\{\bar{\cdot}, \cap, \cup\}$, which is subsequently extended into a model of linear logic under the rules $\{\bar{\cdot}, \cap, \cup, \otimes, \bar{\cdot} \rightarrow \cdot\}$. The introduced prec fits naturally as a multiplicative rule in-between $\otimes$ and $\bar{\cdot} \rightarrow \cdot$, which moreover has the property of commuting with the negation.

Finally, we rederived the proof of equivalence between combs based on quantum states, quantum combs, and networks based on quantum channels [5] as an example of the use of the projective characterization. This equivalence explains in particular why these structures feature a definite signaling direction, which comes with a clear physical implementation. One can wonder how often two differently defined structures with a common no signaling subset are equivalent. And how often do the structures based on the transformation connector (comb-like, which can present ICO in general) reduce in normal form to structures based on a prec (network-like, which cannot present ICO)? Preliminary results indicate these equivalences to be rare, but the question of classifying the isomorphic constructions is left open for future work. Another related but more general question that can be asked in terms of projectors is to determine for which base state structures a given construction based on the transformation reduces to a construction based on the one-way signaling composition. In other words, given an abstract way of constructing higher-order transformations, for which base

state structures does it reduce to a network?

Several aspects of the theory of higher-order transformations have not been explored here and remain open for future work. The composition of types in the sense of Ref. [5] has only been partially addressed. Knowing the projector characterizing a transformation is sufficient to know what will be the projector characterizing its set of outputs. However, if the input state structure is now restricted to a subset with a different projector, is there a rule to apply to the projector of the transformation in order to get the projector of the possibly now restricted output set? For example, an n -comb takes $(n-1)$ -combs as inputs and outputs a 1-comb. One can ask what happens when the inputs are restricted to the no signaling subset: when $(n-1)$ 1-combs are plugged into an n -comb instead, is the output set still the full set of 1-combs?

The notion of causal separability [8, 28, 29, 11] still is to be defined for general higher-order transformations. The link between the prec and single signaling direction, as well as the normal form, should give a basis for the definition: the normal form can be used to decompose any process into a sum of causally ordered pieces. If this sum is convex, this gives a sufficient condition for causal separability. Nonetheless, there remains to tackle the question of dynamical causal orders, which is a probabilistic notion, and therefore not naturally captured by the formalism of projectors.

The issue of realizability has not been explored. That is, given a state structure, is it possible to realize each of its elements in a lab experiment? Is there a systematic way to relate these abstract mathematical objects to a circuit realization, as is the case for quantum combs [3] and for some process matrices in terms of time-delocalized subsystems [30, 31]? One may expect that the normal decomposition of general state structures into the union of several one-way signaling structures will prove useful for answering this question.

Note added: After completion of the main results of this work and during the preparation of the manuscript, we became aware of an independent work by Simmons and Kissinger [26], which also investigates the signaling structure of higher-order quantum transformations. Some of the results presented here were also found in this other work using a different formalism.

Acknowledgments

The authors are grateful for the extensive reviews provided by anonymous reviewers, which contributed greatly to improving this work. They also acknowledge useful technical discussions with Jessica Bavaresco and Aleks Kissinger.

T. H. is grateful to the organizers and participants of the 2021 Sejny Summer Institute for listening to his rambling interventions about nesting supermaps – in particular to Pablo Arnault, Titouan Carette, Natália Móller, Eleftherios Tselentis, and Augustin Vanrietvelde for technical discussions. T. H. is especially indebted to Titouan Carette for pointing out the connection with linear logic. In addition, T.H. would like to thank Esteban Castro-Ruiz and Joseph Cunningham for their help and comments, as well as Alexandra Elbakyan for providing access to the scientific literature. T. H. benefited from the support of the French Community of Belgium within the framework of the financing of a FRIA grant, then from the Slovak grants ID# APVV-22-0570 and ID# VEGA 2/0128/24.

This publication was made possible through the support of the ID# 61466 grant and ID# 62312 grant from the John Templeton Foundation, as part of the ‘The Quantum Information Structure of Spacetime’ Project (QISS), and by the ID# 63683 grant from the John Templeton Foundation, as part of the “Without SpaceTime” Project (WOST).. The opinions expressed in this publication are those of the authors and do not necessarily reflect the views of the John Templeton Foundation. This work was supported by the Program of Concerted Research Actions (ARC) of the Université libre de Bruxelles and from the F.R.S.-FNRS under project CHEQS within the Excellence of Science (EOS) program. O. O. is a Senior Research Associate of the Fonds de la Recherche Scientifique (F.R.S.-FNRS).

Illustrations were drawn using draw.io.

References

- [1] Karl Kraus. “States, effects, and operations: Fundamental notions of quantum theory”. [Volume 190 of Lecture notes in physics](#). Springer-Verlag Berlin Heidelberg. (1983). 1 edition.
- [2] G. Chiribella, G. M. D’Ariano, and

- P. Perinotti. “Transforming quantum operations: Quantum supermaps”. *EPL (Europhysics Letters)* **83**, 30004 (2008). [arXiv:0804.0180](#).
- [3] Giulio Chiribella, Giacomo Mauro D’Ariano, and Paolo Perinotti. “Theoretical framework for quantum networks”. *Physical Review A* **80** (2009). [arXiv:0904.4483](#).
- [4] Paolo Perinotti. “Causal structures and the classification of higher order quantum computations”. *Tutorials, Schools, and Workshops in the Mathematical Sciences* **Page 103–127** (2017). [arXiv:1612.05099](#).
- [5] Alessandro Bisio and Paolo Perinotti. “Theoretical framework for higher-order quantum theory”. *Proceedings of the Royal Society A: Mathematical, Physical and Engineering Sciences* **475**, 20180706 (2019). [arXiv:1806.09554](#).
- [6] Giulio Chiribella, Giacomo Mauro D’Ariano, Paolo Perinotti, and Benoit Valiron. “Quantum computations without definite causal structure”. *Physical Review A* **88**, 022318 (2013). [arXiv:0912.0195](#).
- [7] E. B. Davies and J. T. Lewis. “An operational approach to quantum probability”. *Communications in Mathematical Physics* **17**, 239–260 (1970).
- [8] Ognian Oreshkov, Fabio Costa, and Časlav Brukner. “Quantum correlations with no causal order”. *Nature Communications* **3**, 1–13 (2012). [arXiv:1105.4464](#).
- [9] Aleks Kissinger and Sander Uijlen. “A categorical semantics for causal structure”. *Logical Methods in Computer Science* **Volume 15, Issue 3** (2019). [arXiv:1701.04732](#).
- [10] Mateus Araújo, Cyril Branciard, Fabio Costa, Adrien Feix, Christina Giarmatzi, and Časlav Brukner. “Witnessing causal nonseparability”. *New Journal of Physics* **17**, 102001 (2015). [arXiv:1506.03776](#).
- [11] Timothée Hoffreumon and Ognian Oreshkov. “The multi-round process matrix”. *Quantum* **5**, 384 (2021). [arXiv:2005.04204](#).
- [12] Simon Milz, Jessica Bavaresco, and Giulio Chiribella. “Resource theory of causal connection”. *Quantum* **6**, 788 (2022). [arXiv:2110.03233](#).
- [13] Simon Milz and Marco Túlio Quintino. “Characterising transformations between quantum objects, of quantum properties, and transformations without a fixed causal order”. *Quantum* **8**, 1415 (2024). [arXiv:2305.01247](#).
- [14] Jean-Yves Girard. “Linear logic”. *Theoretical Computer Science* **50**, 1–101 (1987).
- [15] Andrej Jamiołkowski. “Linear transformations which preserve trace and positive semidefiniteness of operators”. *Reports on Mathematical Physics* **3**, 275–278 (1972).
- [16] Man-Duen Choi. “Positive linear maps on complex matrices”. *Linear Algebra and its Applications* **10**, 285–290 (1975).
- [17] David Beckman, Daniel Gottesman, Michael A. Nielsen, and John Preskill. “Causal and localizable quantum operations”. *Physical Review A* **64** (2001). [arXiv:quant-ph/0102043](#).
- [18] Marco Piani, Michał Horodecki, Paweł Horodecki, and Ryszard Horodecki. “Properties of quantum nonsignaling boxes”. *Physical Review A* **74** (2006). [arXiv:quant-ph/0505110](#).
- [19] Michael A. Nielsen and Isaac L. Chuang. “Quantum computation and quantum information”. *Cambridge University Press*. (2009). 2 edition.
- [20] Giulio Chiribella, Giacomo M. D’Ariano, and Paolo Perinotti. “Memory effects in quantum channel discrimination”. *Physical Review Lett.* **101**, 180501 (2008). [arXiv:0803.3237](#).
- [21] Sally Shrapnel, Fabio Costa, and Gerard Milburn. “Updating the Born rule”. *New Journal of Physics* **20** (2018). [arXiv:1702.01845](#).
- [22] Man-Duen Choi and Edward G. Effros. “Injectivity and operator spaces”. *Journal of Functional Analysis* **24**, 156 – 209 (1977).
- [23] Fumio Hiai and Dénes Petz. “Introduction to matrix analysis and applications”. *Universitext*. Springer, Cham. (2014).
- [24] Esteban Castro-Ruiz, Flaminia Giacomini, and Časlav Brukner. “Dynamics of Quantum Causal Structures”. *Physical Review X* **8**, 1–18 (2018). [arXiv:1710.03139](#).
- [25] Esteban Castro-Ruiz (2019). private communication.
- [26] Will Simmons and Aleks Kissinger. “Higher-order causal theories are models of bv-logic”. In Stefan Szeider, Robert Ganian, and Alexandra Silva, editors, 47th Interna-

- tional Symposium on Mathematical Foundations of Computer Science (MFCS 2022). Volume 241 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pages 80:1–80:14. Dagstuhl, Germany (2022). Schloss Dagstuhl – Leibniz-Zentrum für Informatik. [arXiv:2205.11219](#).
- [27] Giulio Chiribella, Giacomo Mauro D’Ariano, and Paolo Perinotti. “Informational derivation of quantum theory”. *Physical Review A* **84** (2011). [arXiv:1011.6451](#).
- [28] Ognian Oreshkov and Christina Giarmatzi. “Causal and causally separable processes”. *New Journal of Physics* **18**, 093020 (2016). [arXiv:1506.05449](#).
- [29] Julian Wechs, Alastair A Abbott, and Cyril Branciard. “On the definition and characterisation of multipartite causal (non)separability”. *New Journal of Physics* **21**, 013027 (2019). [arXiv:1807.10557](#).
- [30] Ognian Oreshkov. “Time-delocalized quantum subsystems and operations: on the existence of processes with indefinite causal structure in quantum mechanics”. *Quantum* **3**, 206 (2019). [arXiv:1801.07594](#).
- [31] Julian Wechs, Cyril Branciard, and Ognian Oreshkov. “Existence of processes violating causal inequalities on time-delocalised subsystems”. *Nature Communications* **14**, 1471 (2023). [arXiv:2201.11832](#).
- [32] Steven Roman. “Advanced linear algebra”. *Springer*. (2008). 3 edition.
- [33] Robert Piziak, Patrick L. Odell, and R. Hahn. “Constructing projections on sums and intersections”. *Computers & Mathematics with Applications* **37**, 67–74 (1999).
- [34] Tomoyuki Morimae. “The process matrix framework for a single-party system” (2014). [arXiv:1408.1464](#).

A Superoperator projectors

This appendix explores the properties of the orthogonal superoperator projectors that are used as the principal tool for the characterization in the main text. These are defined under the notion of a projector on an operator system in Section A.1. Six operations whose purpose is motivated by the main text are then defined on this set. These are: the intersection and union, $\cap, \cup$,

the negation, $\neg$, the tensor and transformation, $\otimes, \rightarrow$, and the prec, $\prec$, successively reviewed in Sec. A.2, A.3, A.4, and A.5. Some important properties of these operations are proven for the purposes of the main text, among which the fact that each operation applied on commuting projector(s) on operator system(s) results in commuting projector(s) on operator system(s), justifying the assumption of commutation while proving Theorem 1.

A.1 Definition and required properties of superoperator projectors

Projectors on operator systems. As explained below Def. 1, a projector on an operator system $\mathcal{P}_A : \mathcal{L}(\mathcal{H}^A) \rightarrow \mathcal{L}(\mathcal{H}^A)$ is a linear orthogonal superoperator projector to a self-adjoint subspace that contains the identity.

A linear superoperator $\mathcal{P}_A$ is a *projector* on a subspace of $\mathcal{L}(\mathcal{H}^A)$ if and only if it is idempotent,

$$\mathcal{P}_A \circ \mathcal{P}_A = \mathcal{P}_A. \quad (78)$$

An operator $V \in \mathcal{L}(\mathcal{H}^A)$ belongs to the subspace defined by $\mathcal{P}_A$ if and only if

$$\mathcal{P}_A \{V\} = V. \quad (79)$$

$\mathcal{P}_A$ projects to a subspace closed under the adjoint if and only if it obeys

$$\mathcal{P}_A \{V\} = V \Rightarrow \mathcal{P}_A \{V^\dagger\} = V^\dagger, \quad (80)$$

which can be written concisely as $\mathcal{P}_A \circ \dagger = \dagger \circ \mathcal{P}_A$ where $\dagger$ means ‘taking the adjoint in $\mathcal{L}(\mathcal{H}^A)$ ’.

A projector is *orthogonal* if it does not increase the norm of operators:

$$\|\mathcal{P}_A \{V\}\| \leq \|V\|, \quad (81)$$

where $\|V\| \equiv \sqrt{\text{Tr}[V^\dagger \cdot V]}$ is the Hilbert-Schmidt norm. Note that this condition is equivalent to requiring the projector to be self-adjoint with respect to the Hilbert-Schmidt inner product $\langle V', V \rangle \equiv \text{Tr}[V'^\dagger \cdot V]$ (see *e.g.* Ref. [32, Theorem 10.5]),

$$\text{Tr}[\mathcal{P}_A \{V'\}^\dagger \cdot V] = \text{Tr}[V'^\dagger \cdot \mathcal{P}_A \{V\}], \quad \forall V, V'. \quad (82)$$

As the projectors appearing in the main text are used to talk about the subspace spanned by a

collection of elements rather than the elements themselves, what is truly relevant in the projective characterization is the image of the projectors, not the projectors themselves. Consequently, there is no loss of generality in assuming every projector to be orthogonal since there always exists an orthogonal projector with the same image as any given projector. Having this property will greatly simplify some of the proofs. For a linear map $\mathcal{M} \in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^A))$, we note its Hilbert-Schmidt adjoint with a $*$ i.e., $\mathcal{M}^*$ is the unique map satisfying $\langle V', \mathcal{M}(V) \rangle = \langle \mathcal{M}^*(V'), V \rangle$ for all V, V' . Using this notation, the self-adjointness condition can be written concisely as $\mathcal{P}_A = \mathcal{P}_A^*$.

As the projectors should project onto operator systems, which are subspaces that contain the identity, it should be true that their images contain the span of the identity as a subspace. A necessary and sufficient condition for that to be true is

$$\mathcal{P}_A \circ \mathcal{D}_A = \mathcal{D}_A \circ \mathcal{P}_A = \mathcal{D}_A, \quad (83)$$

with $\mathcal{D}_A$ defined as in Eq. (11).

Definition and example. Summarizing, a projector on an operator system is a linear super-operator $\mathcal{P}_A \in \mathcal{L}(\mathcal{L}(\mathcal{H}^A), \mathcal{L}(\mathcal{H}^A))$ obeying the defining conditions (78), (80), and (83). Moreover, for the purpose of this article, it is assumed to be self-adjoint without loss of generality, condition (82).

Putting everything together, the projector appearing in Eq. (8c) is an orthogonal projector on an operator system, meaning it obeys the following set of conditions:

$$\mathcal{P}_A \circ \mathcal{P}_A = \mathcal{P}_A; \quad (84a)$$

$$\mathcal{P}_A \circ \dagger = \dagger \circ \mathcal{P}_A; \quad (84b)$$

$$\mathcal{P}_A^* = \mathcal{P}_A; \quad (84c)$$

$$\mathcal{P}_A \circ \mathcal{D}_A = \mathcal{D}_A \circ \mathcal{P}_A = \mathcal{D}_A. \quad (84d)$$

When referring to ‘a projector’ in this article, we implicitly mean ‘an orthogonal projector on an operator system’, i.e., ‘a projector obeying conditions (84)’, unless said otherwise.

Examples of projectors on operator systems over a 2-dimensional Hilbert space can be generated using the Pauli basis $\{\mathbb{1}, \sigma_x, \sigma_y, \sigma_z\}$ (see e.g., Ref. [19] for its definition and properties): it can

be checked that any projector ‘removing’ one or several basis elements other than the identity is a projector on an operator system i.e., obeys (84). Such is the case of the projector to the subspace of operators that are diagonal in the computational basis, which is used for quantum decoherence. This projector, noted $\Delta : \mathcal{L}(\mathcal{H}^A) \rightarrow \mathcal{L}(\mathcal{H}^A)$ and which acts as $\Delta(\cdot) \equiv \frac{1}{2} \text{Tr}[\cdot] + \frac{\sigma_z}{2} \text{Tr}[\sigma_z \cdot]$, is the projector to the subspace spanned by $\{\mathbb{1}, \sigma_z\}$, thus a projector on an operator system.

Sets of commuting base projectors. For the purposes of this work, it is assumed that all the projectors under consideration always commute pairwise when acting on the same tensor factor of a Hilbert space. Remark that the commutation of two projectors is a sufficient condition for their composition to be a projector as well (see e.g., Ref. [32, Theorem 2.26 3]).

At first glance, the physical heuristic behind the assumption is that we are only interested in building (higher-order) maps between subsystems of a fixed kind: if Alice’s system is a qubit and Bob’s is a classical bit, this will not change whether we consider their joint state or a channel between them; the framework developed in this article is only used to characterize their joint state space or the set of channels they can share, not what happens when Alice suddenly demotes her system to a bit. The characterization is done by combining the base projectors of Alice and Bob, which in this case are $\mathcal{I}_A$ and Δ_B (where $\mathcal{I}$ is the identity projector and Δ is the projector to the diagonal subspace as defined above). For that reason, it can be assumed that there is only one base projector associated with each factor of a Hilbert space at the start of the construction.

The rules introduced in the main text and detailed below then lead to several ways to combine the projectors of Alice and Bob into a projector acting on space $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$. Hence, the ‘Alice+Bob’ system can be of several different natures: it can be bipartite states, but also channels (in CJ representation), bipartite measurements, etc. All of these correspond to different projectors, and each of these projectors can in turn be seen as the base projectors in the coarse-grained, ‘Alice+Bob’ description. One of the things shown in this appendix is precisely that every introduced rule to build new projectors out of known ones will preserve commutation. Hence, any two pro-

jectors characterizing systems of a different type and obtained by relating the same base projectors on the same subsystems using any of the various algebraic rules defined in the following will commute by construction.

Non-singleton set of base projectors.

Upon closer inspection, it may still be interesting to study what happens when the base subsystems can be picked among a set of possible types. For instance, it is very natural to extrapolate the characterization of a class of processes to a case in which, say, Alice's state preparation becomes a discarding procedure (graphically, it amounts to changing Alice's wire from facing up to facing down). This replacement actually corresponds to the 'not' operation in the algebra, so it will preserve commutation as will be proven below. Hence, any base projector can be associated with a second base projector, its negation. In addition, the operation of 'intersection' and 'union' will also result in two new projectors that commute with both the base projector and its negation. Therefore, even when one considers one base projector per subsystem, it is always extensible to a set of at least four pairwise commuting base projectors. There is no point in restricting the set of base projectors then and one can consider any commuting set of base projectors in general. Therefore, for the projective characterization developed in this work, the general setting considers k parties $A, B, \dots$, each of which having access to $n_A, n_B, \dots$ base projectors $\{\mathcal{P}_A^{(1)}, \mathcal{P}_A^{(2)}, \dots, \mathcal{P}_A^{(n_A)}\}$, $\{\mathcal{P}_B^{(1)}, \mathcal{P}_B^{(2)}, \dots, \mathcal{P}_B^{(n_B)}\}$, ... corresponding to the different types of subsystems they may have access to. The commutation assumption, which can be expressed as

$$\mathcal{P}_X^{(i)} \circ \mathcal{P}_X^{(j)} = \mathcal{P}_X^{(j)} \circ \mathcal{P}_X^{(i)}, \quad \forall X, \forall i, j. \quad (85)$$

As mentioned in the related comment made at the end of Sec. 5.4, these will not only encompass all higher-order objects that could have been built using a fine-graining of the subsystem. For example, it can also include scenarios in which Alice can choose whether she wants to use a classical or quantum bit as her system to model scenarios featuring potential decoherence. In our framework, this corresponds to Alice having a set of base projectors made of two elements, $\{\mathcal{I}_A, \Delta_A\}$ (which will then be extended to a set of eight elements by using the not, intersection and union

operations). We want to bring to the reader's attention that since these two projectors commute, all the results of this paper will still hold, albeit it is impossible to build Δ out of $\mathcal{I}$ using any of the rules presented in this article. This is thus a situation characterized by our result but not considered in this article, which opens a potential direction for future research.

Commutation with the transposition.

A follow-up consequence of assuming that all projectors commute pairwise is that, when acting on self-adjoint operators, they can be assumed to commute with the transposition induced by the CJ isomorphism. This condition will be required simply because it removes the necessity of keeping track of the transposes, resulting in less clustered formulae.

Here, transposition is meant in a more general sense than matrix transpose. The CJ correspondence depends on a choice of basis, namely the one in which the matrix transpose that appears in Eq. (1) is performed. This implicitly means that the standard bases of $\mathcal{H}^A$ and $\mathcal{H}^B$ were chosen to define the isomorphism. However, this is an arbitrary choice. Eq. (1) is a shortcut for

$$M = \text{Tr}_{A'B'}[(\phi_{AA'}^+ \otimes \phi_{B'B}^+)(\mathbb{1}_A \otimes (\mathcal{I} \otimes \mathcal{M})\{\phi^+\}_{A'B'} \otimes \mathbb{1}_B)], \quad (86)$$

and it is only because ϕ^+ was chosen to be $\phi^+ = \sum_{i,j} |i\rangle \langle j| \otimes |i\rangle \langle j|$ that it corresponds to matrix transposition in the standard basis. Yet, any other choice of unnormalized maximally entangled density matrix Φ is equally good. Let $\{e_\mu\}_{\mu=0}^{d_A^2-1}$ be an orthonormal basis of $\mathcal{L}(\mathcal{H}^A)$. Then this basis can always be associated with an unnormalized maximally entangled state of the form $\Phi_{A'A} = \sum_\mu e_\mu \otimes e_\mu$, which is in turn used to define a transposition with respect to $\{e_\mu\}$ by

$$V_A^T \equiv \text{Tr}_{A'}[\Phi(V_{A'} \otimes \mathbb{1}_A)], \quad \forall V_{A'} \in \mathcal{L}(\mathcal{H}^{A'}). \quad (87)$$

This abstract notion of transposition happens to correspond to the one of matrix transpose with respect to a basis $\{|\varphi_i\rangle\}_{i=0}^{d_A-1}$ of $\mathcal{H}^A$ precisely when the basis $\{e_\mu\}$ is $\{e_\mu := |\varphi_i\rangle \langle \varphi_j|_{\mu=(i,j)}\}$, but it does not have to have this form to define a CJ correspondence. All that matters is that a pair of bases for $\mathcal{L}(\mathcal{H}^A)$ and $\mathcal{L}(\mathcal{H}^B)$ is chosen.

This freedom of basis in the definition of the CJ correspondence can be leveraged to make all

projectors commute with the transpositions it induces. For instance, this work typically deals with equations of the form

$$\mathcal{M}(V_A) = \mathcal{P}_B \{ \mathcal{M}(\mathcal{P}_A \{V_A\}) \}, \quad (88)$$

where $\mathcal{M}$ is CP map, V_A a positive operator, and $\mathcal{P}_A$ and $\mathcal{P}_B$ are projectors on operator systems. Our methods will use the reverse direction of the CJ correspondence (2) to rephrase the above into

$$\mathcal{M}(V_A) = \mathcal{P}_B \left\{ \left(\text{Tr}_A [M(\mathcal{P}_A \{V_A\} \otimes \mathbb{1})] \right)^T \right\}, \quad (89)$$

and then gather all projectors into a single expression as in

$$\mathcal{M}(V_A) = \left(\text{Tr}_A [((\mathcal{P}_A \otimes \mathcal{P}_B) \{M\}) (V_A \otimes \mathbb{1})] \right)^T, \quad (90)$$

where the tensor product of projectors is defined in a standard way (see Eq. (19) or in App. A.4.1 below). Going from Eq. (89) to (90) assumed that $(\mathcal{P}_B \{N_B\})^T = \mathcal{P}_B \{N_B^T\}$ for all N_B . A similar rewriting can also be made when computing $\mathcal{M}^*(N_B)$, in which case one needs to assume that $(\mathcal{P}_A \{V_A\})^T = \mathcal{P}_A \{V_A^T\}$ for all V_A . When all operators are self-adjoint, we now show that these assumptions can always be satisfied by picking a specific choice of bases to define the CJ isomorphism.

The general statement we prove is that if the base projectors form a set of n_X pairwise commuting orthogonal projectors on operator systems over subsystem X , $\{\mathcal{P}_X^{(1)}, \mathcal{P}_X^{(2)}, \mathcal{P}_X^{(3)}, \dots, \mathcal{P}_X^{(n_X)}\}$, the following can be ensured:

$$(\mathcal{P}_X^{(i)} \{V\})^T = \mathcal{P}_X^{(i)} \{V^T\}, \quad \forall i, \forall V : V = V^\dagger, \quad (91)$$

for the transpositions induced by the CJ isomorphism. This is done by choosing the basis for the isomorphism to be the one diagonalizing the set of projectors.

If this basis is chosen to be the one that diagonalizes the set of base projectors, i.e. if the base projectors is a set of n pairwise commuting orthogonal projectors $\{\mathcal{P}_X^{(1)}, \mathcal{P}_X^{(2)}, \mathcal{P}_X^{(3)}, \dots, \mathcal{P}_X^{(n_X)}\}$ expressed in this basis as $\mathcal{P}_X^{(i)} \{\cdot\} = \sum_\mu c_\mu^{(i)} e_\mu \text{Tr} [e_\mu^\dagger \cdot]$ with $i = 1, \dots, n_X$ and $c_\mu^{(i)} \in \{0, 1\}$, then

$$\sum_\nu \mathcal{P}_{X'}^{(i)} \{e_\nu\} \otimes e_\nu = \sum_\nu e_\nu \otimes \mathcal{P}_X^{(i)} \{e_\nu\}, \quad (92)$$

and it follows that the projectors commute with the transpose when applied on self-adjoint operators. Indeed, $\forall V : V = V^\dagger$ and $\forall i$:

$$\begin{aligned} & \mathcal{P}_X^{(i)} \{V^T\} \\ &= \mathcal{P}_A^{(i)} \left\{ \text{Tr}_{X'} \left[(V_{X'} \otimes \mathbb{1}_X) \left(\sum_\nu e_\nu \otimes e_\nu \right) \right] \right\} \\ &= \text{Tr}_{X'} \left[(V_{X'} \otimes \mathbb{1}_X) \left(\sum_\nu e_\nu \otimes \mathcal{P}_X^{(i)} \{e_\nu\} \right) \right] \\ &= \text{Tr}_{X'} \left[(V_{X'} \otimes \mathbb{1}_X) \left(\sum_\nu \mathcal{P}_{X'}^{(i)} \{e_\nu\} \otimes e_\nu \right) \right] \\ &= \text{Tr}_{A'} \left[\left(\mathcal{P}_{X'}^{(i)} \{V_{X'}\} \otimes \mathbb{1}_X \right) \left(\sum_\nu e_\nu \otimes e_\nu \right) \right] \\ &= \left(\mathcal{P}_X^{(i)} \{V\} \right)^T, \end{aligned} \quad (93)$$

where passing from the third to the fourth line was possible only because $V = V^\dagger$ and $\mathcal{P}_X^{(i)}$ is self-adjoint by definition.

Definition. Summarizing, any subsystem X will be associated with n_X different types given a priori, called their base types. These base types correspond to n_X different state structures to which are associated n_X base projectors noted $\{\mathcal{P}_X^{(1)}, \mathcal{P}_X^{(2)}, \mathcal{P}_X^{(3)}, \dots, \mathcal{P}_X^{(n_X)}\}$ or $\{\mathcal{P}_X, \mathcal{P}'_X, \mathcal{P}''_X, \dots\}$. For every self-adjoint operator $V \in \mathcal{L}(\mathcal{H}^X)$, the following is always assumed to hold for a set of base projectors:

$$\mathcal{P}_X^{(i)} \circ \mathcal{P}_X^{(j)} = \mathcal{P}_X^{(j)} \circ \mathcal{P}_X^{(i)}, \quad \forall i, j; \quad (94a)$$

$$(\mathcal{P}_X^{(i)} \{V\})^T = \mathcal{P}_X^{(i)} \{V^T\}, \quad \forall i. \quad (94b)$$

Remark: non-pairwise-commuting set of base projectors. The assumption that the base projectors all commute is reasonable for the scope of this paper and potential extensions, as argued in the above discussion. Yet it is not necessary: it was only made because there are no heuristical reasons to go beyond it. While it greatly simplifies the mathematics by allowing the intersection (or cap) operation $\cap$ to be defined as the composition of two projectors (Eq. (25) in the main text; its mathematical properties are studied in the next subsection below), dropping it may also open a direction for further research: the natural follow-up question to considering a higher-order theory based on a set of pairwise commuting base projectors would be to

study what happens when the base projectors are no longer commuting. In this case, most of the results would have to be generalized in a version that does not use this assumption nor the related one that the projectors commute with the transpose. Nonetheless, besides pure mathematical curiosity, we do not have a clear idea of what the purpose of such a framework may be; let alone its physical interpretation.

A.2 Adding cap and cup operations makes an algebra

Here, we prove that superoperator projectors on operator systems constitute an algebra under certain rules. Assume $\mathcal{P}_A$ and $\mathcal{P}'_A$ to be two arbitrary projectors on operator systems, not necessarily the same. By definition, a projector is $\mathbb{C}$ -linear:

$$\mathcal{P}_A \left\{ \sum_i q_i V_i \right\} = \sum_i q_i \mathcal{P}_A \{V_i\} , \quad (95)$$

$q_i \in \mathbb{C}$. This linearity allows us to carry on the addition ‘+’ of $\mathcal{L}(\mathcal{H}^A)$ at the level of the projectors:

$$(\mathcal{P}_A + \mathcal{P}'_A) \{V\} \equiv \mathcal{P}_A \{V\} + \mathcal{P}'_A \{V\} . \quad (96)$$

Using linearity again, one can define the negation ‘-’ and scalar multiplication of projectors, thereby defining a vector space over $\mathbb{C}$. Associativity and commutativity are not hard to prove from there.

Projectors also have a natural ‘conjunction’ operation that can be interpreted as a multiplication or as a logic ‘and’, here nicknamed the **intersection** or the **cap** and noted via the ‘ $\cap$ ’ symbol. In the case of commuting projectors $\mathcal{P}_A$ and $\mathcal{P}'_A$, it is

$$\mathcal{P}_A \cap \mathcal{P}'_A \equiv \mathcal{P}_A \circ \mathcal{P}'_A . \quad (97)$$

This projector $\mathcal{P}_A \cap \mathcal{P}'_A$ characterizes the intersection of the operator systems $\mathcal{A}$ and $\mathcal{A}'$, see the main text. It is straightforward to prove that it fits the definition of a conjunction: it is a binary operation distributive under the addition defined above and compatible with the scalar multiplication; it is both associative and commutative; for projectors satisfying Eqs. (94), it is resulting in a projector satisfying Eqs. (94); and it is idempotent when acting on projectors since:

$$\mathcal{P}_A \cap \mathcal{P}_A \equiv \mathcal{P}_A \circ \mathcal{P}_A \equiv \mathcal{P}_A^2 = \mathcal{P}_A . \quad (98)$$

In the above equation, we have also defined a shorthand notation for ‘squaring’ under the multiplication $\cap$: $\mathcal{P}_A \cap \mathcal{P}_A \equiv \mathcal{P}_A^2$.

An issue revealed by squaring is that the operation ‘+’ does not necessarily map projectors to a projector:

$$(\mathcal{P}_A + \mathcal{P}'_A)^2 = \mathcal{P}_A^2 + \mathcal{P}'_A^2 + 2(\mathcal{P}_A \cap \mathcal{P}'_A) . \quad (99)$$

As the last term of the right-hand side is not always zero, $(\mathcal{P}_A + \mathcal{P}'_A)^2 \neq \mathcal{P}_A + \mathcal{P}'_A$ in general, so addition does not preserve idempotency. We wish nonetheless to have operations that keep us in the vector space of commuting projectors, so we redefine the addition to be the union of two projectors (see *e.g.* Ref. [33]). This is inspired by set theory: as $\mathcal{P}_A$ maps to a subspace $\mathcal{A} \subset \mathcal{L}(\mathcal{H}^A)$ and $\mathcal{P}'_A$ to $\mathcal{A}' \subset \mathcal{L}(\mathcal{H}^A)$ we want an addition ‘ $\cup$ ’ so that $\mathcal{P}_A \cup \mathcal{P}'_A$ maps to $\mathcal{A} \cup \mathcal{A}'$. This requirement is realized by the ‘disjunction’ operation, that can be interpreted as an addition or as logic ‘or’, nicknamed the **union** or the **cup** (see also Ref. [11]):

$$\mathcal{P}_A \cup \mathcal{P}'_A \equiv \mathcal{P}_A + \mathcal{P}'_A - \mathcal{P}_A \cap \mathcal{P}'_A . \quad (100)$$

As the name hints, $\mathcal{P}_A \cup \mathcal{P}'_A$ characterizes the union of the operator systems $\mathcal{A}$ and $\mathcal{A}'$. It is again straightforward to prove that it fits the definition from the properties inherited by the addition and the cap: the cup is a binary operation distributive under the addition defined above and compatible with the scalar multiplication; it is both associative and commutative; for projectors satisfying Eqs. (94), it results in a projector satisfying Eqs. (94); and it inherits the idempotency from the idempotency of the cap as $\mathcal{P}_A \cup \mathcal{P}_A = \mathcal{P}_A + \mathcal{P}_A - (\mathcal{P}_A \cap \mathcal{P}_A) = \mathcal{P}_A + \mathcal{P}_A - \mathcal{P}_A$.

These caps and cup connectors give a quick way to prove subspace inclusions. The subspace spanned by state structure $\mathcal{A}'$ is embedded in the subspace spanned by another state structure $\mathcal{A}$ if and only if either of the following holds (see *e.g.*, Ref. [32])

$$\mathcal{P}_A \cap \mathcal{P}'_A = \mathcal{P}'_A ; \quad (101a)$$

$$\mathcal{P}_A \cup \mathcal{P}'_A = \mathcal{P}_A . \quad (101b)$$

The intersection and union of two projectors project to, respectively, the intersection and union of their images [32, Theorem 2.26] *i.e.*,

$$\text{Im} \{ \mathcal{P}_A \cap \mathcal{P}'_A \} = \text{Im} \{ \mathcal{P}_A \} \cap \text{Im} \{ \mathcal{P}'_A \} , \quad (102a)$$

$$\text{Im} \{ \mathcal{P}_A \cup \mathcal{P}'_A \} = \text{Im} \{ \mathcal{P}_A \} \cup \text{Im} \{ \mathcal{P}'_A \} , \quad (102b)$$

where the $\cap$ and $\cup$ appearing in the right-hand side of the above are the set-theoretic notions of intersection and union, respectively. Condition (84d) is then recast into

$$\mathcal{P}_A \cap \mathcal{D}_A = \mathcal{D}_A. \quad (103)$$

As defined in Eq. (28) of the main text, obeying Eqs. (101) is concisely noted as $\mathcal{P}'_A \subseteq \mathcal{P}_A$ and means that such an equation must be understood in terms of the image of the projectors. Accordingly, an equivalence of two projectors is defined as the equivalence of their images:

$$\mathcal{P}_A = \mathcal{P}'_A \iff \text{Im}\{\mathcal{P}_A\} = \text{Im}\{\mathcal{P}'_A\}. \quad (104)$$

The algebra. Commuting projectors on operator systems thus form an algebra over $\mathbb{C}$ since they are closed under $(\cap, \cup)$ operations as well as scalar multiplication. The cap and cup can indeed be seen as some multiplication and addition of projectors as the former distributes over the latter:

$$(\mathcal{P}_A \cup \mathcal{P}'_A) \cap \mathcal{P}''_A = (\mathcal{P}_A \cap \mathcal{P}''_A) \cup (\mathcal{P}'_A \cap \mathcal{P}''_A). \quad (105)$$

Indeed, $(\mathcal{P}_A \cup \mathcal{P}'_A) \cap \mathcal{P}''_A = (\mathcal{P}_A + \mathcal{P}'_A - (\mathcal{P}_A \cap \mathcal{P}'_A)) \cap \mathcal{P}''_A = (\mathcal{P}_A \cap \mathcal{P}''_A) + (\mathcal{P}'_A \cap \mathcal{P}''_A) - (\mathcal{P}_A \cap \mathcal{P}'_A) \cap \mathcal{P}''_A$, and $(\mathcal{P}_A \cap \mathcal{P}'_A) \cap \mathcal{P}''_A = (\mathcal{P}_A \cap \mathcal{P}'_A) \cap (\mathcal{P}''_A \cap \mathcal{P}''_A) = (\mathcal{P}_A \cap \mathcal{P}'_A) \cap (\mathcal{P}''_A \cap \mathcal{P}''_A) = (\mathcal{P}_A \cap \mathcal{P}'_A) \cap (\mathcal{P}''_A \cap \mathcal{P}''_A)$.

The projectors $\mathcal{I}$ and $\mathcal{D}$ defined in Eqs. (9) and (11) above play a special role in it: one can identify the identity superoperator $\mathcal{I}$ as the multiplicative identity or ‘**unit**’ of the algebra of projectors since for all projectors $\mathcal{P}$,

$$\mathcal{I} \cap \mathcal{P} = \mathcal{P}; \quad (106)$$

it is also the additive absorbing element i.e.

$$\mathcal{I} \cup \mathcal{P} = \mathcal{I}. \quad (107)$$

The other way around, the depolarizing superoperator $\mathcal{D}$ is the multiplicative absorbing element or ‘**zero**’ of the algebra,

$$\mathcal{D} \cap \mathcal{P} = \mathcal{D}, \quad (108)$$

and the additive identity,

$$\mathcal{D} \cup \mathcal{P} = \mathcal{P}. \quad (109)$$

Because of that property, it should be clear that for any element of the algebra, the following is true:

$$\mathcal{D} \subseteq \mathcal{P} \subseteq \mathcal{I}. \quad (110)$$

As a consequence, this is an algebra of idempotent elements equipped with a partial ordering $\subseteq$ that have a common greatest element $\mathcal{I}$ and a common least element $\mathcal{D}$. Moreover, the intersection and union of any two elements $\mathcal{P}, \mathcal{P}'$ are uniquely defined elements of the algebra. It is direct to check that the following property holds

$$\mathcal{P} \cap \mathcal{P}' \subseteq \mathcal{P} \subseteq \mathcal{P} \cup \mathcal{P}', \quad (111)$$

and so does the analog with $\mathcal{P}'$. In other words, the algebra has a lattice structure (which directly comes from the lattice structure of subspace inclusions; see e.g., Ref. [32]).

Remark. Although referring to the cap and cup as ‘multiplication’ and ‘addition’ provides an intuitive meaning to the connectors, in the following we will call them respectively ‘additive disjunction and conjunction’. This choice is made in order to avoid ambiguity with the multiplicative conjunction and disjunction that will be introduced in a next section.

A.3 Adding negation makes a Boolean algebra

In the algebra of superoperator projectors, the new object that naturally appears in Proposition 1, $\overline{\mathcal{P}}_A$, can be seen as an operation on the original projector $\mathcal{P}_A$, whence the ‘bar over $\mathcal{P}_A$ ’ notation. This new operation is defined for any projector, and it promotes the algebra to a **Boolean algebra**.

It is a Boolean algebra if, in addition to the conjunction (‘multiplication’) operation $\cap$ and the disjunction (‘addition’) operation $\cup$, any element is idempotent and has a well-defined complement (or logic ‘not’, thereby referred to as ‘negation’¹⁵) characterized by the condition that the addition of any projector with its negation yields the additive identity, i.e.

$$\mathcal{P}_A \cup \overline{\mathcal{P}}_A = \mathcal{I}_A, \quad (112)$$

¹⁵We could also have used the terminology ‘inverse’, ‘dual’, or we could have stuck with ‘complement’, but these words are already used for different concepts in quantum information and functional analysis. This choice is made by consequence to avoid ambiguities.

where $\overline{\mathcal{P}}_A$ is the negation of $\mathcal{P}_A$. This is exactly what the quasi-orthogonal complement of Proposition 1 is doing. First, compute the cap as $\mathcal{P}_A \cap \overline{\mathcal{P}}_A = \mathcal{P}_A \circ (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) = \mathcal{P}_A - \mathcal{P}_A^2 + \mathcal{D}_A = \mathcal{D}_A = \mathcal{P}_A - \mathcal{P}_A + \mathcal{D}_A$, so

$$\mathcal{P}_A \cap \overline{\mathcal{P}}_A = \mathcal{D}_A. \quad (113)$$

Then, it is true for the cup because $\mathcal{P}_A \cup \overline{\mathcal{P}}_A = \mathcal{P}_A + \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A - (\mathcal{P}_A \cap \overline{\mathcal{P}}_A) = \mathcal{I}_A$. Thus, we define the **negation** of any projector $\mathcal{P}_A$ as

$$\overline{\mathcal{P}}_A \equiv \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A. \quad (114)$$

In addition to property (112), one can verify from Eq. (113) that the projector and its negation intersect at the zero of the algebra, which in this case is indeed $\mathcal{D}_A$ (refer to the diagram in Fig. 2). As it should be, negation is an involution:

$$\overline{\overline{\mathcal{P}}_A} = \mathcal{P}_A, \quad (115)$$

and it is closed so that the negation of any projector is a projector as we now prove.

Negation preserves projector on operator system: Note that since $\mathcal{I}_A$ and $\mathcal{D}_A$ are Hermitian-preserving (HP) self-adjoint (SA), and commute with the transpose (CT), Eqs. (84b), (84c), and (94b), the negation of a projector is HP, SA, and CT provided the original projector is. The idempotency property (84a) is preserved, $\overline{\mathcal{P}}_A^2 = (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) \cap (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) = \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A - \mathcal{P}_A + \mathcal{P}_A^2 - \mathcal{D}_A + \mathcal{D}_A - \mathcal{D}_A + \mathcal{D}_A = \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A$, hence

$$\overline{\mathcal{P}}_A^2 = \overline{\mathcal{P}}_A. \quad (116)$$

To prove that the negation of a projector on an operator system is a projector on an operator system itself, it remains to show that the identity is still contained in the negation, Eq. (84d),

$$\mathcal{D}_A \cap \overline{\mathcal{P}}_A = \mathcal{D}_A. \quad (117)$$

This follows from $\mathcal{D}_A \cap \overline{\mathcal{P}}_A = \mathcal{D}_A \cap (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) = \mathcal{D}_A - \mathcal{D}_A + \mathcal{D}_A$. If we did not have the heuristic of Proposition 1, this could be interpreted as a reason for choosing Eq. (114) as the definition of negation instead of the orthogonal complement $\mathcal{I}_A - \mathcal{P}_A$. In this latter case, identity does not belong to the negated state structure.

Negation preserves commutativity: It should also be clear from $\mathcal{P}_A \cap \overline{\mathcal{P}}_A = \mathcal{P}_A \cap (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) = \mathcal{P}_A - \mathcal{P}_A^2 + \mathcal{D}_A = (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) \cap \mathcal{P}_A$ that negated projectors commute with the original ones, and, by the same kind of proof, that the negations of two commuting projectors still commute:

$$\mathcal{P}_A \cap \mathcal{P}'_A = \mathcal{P}'_A \cap \mathcal{P}_A \implies \overline{\mathcal{P}}_A \cap \overline{\mathcal{P}'_A} = \overline{\mathcal{P}'_A} \cap \overline{\mathcal{P}}_A. \quad (118)$$

De Morgan duality: Moreover, the De Morgan laws are valid for the Boolean algebra of projectors:

$$\overline{\mathcal{P}_A \cup \mathcal{P}'_A} = \overline{\mathcal{P}_A} \cap \overline{\mathcal{P}'_A}, \quad (119a)$$

$$\overline{\mathcal{P}_A \cap \mathcal{P}'_A} = \overline{\mathcal{P}_A} \cup \overline{\mathcal{P}'_A}. \quad (119b)$$

The proof is more involved: $\overline{\mathcal{P}_A \cup \mathcal{P}'_A} = \mathcal{I}_A - \mathcal{P}_A \cup \mathcal{P}'_A + \mathcal{D}_A = \mathcal{I}_A - (\mathcal{P}_A + \mathcal{P}'_A - (\mathcal{P}_A \cap \mathcal{P}'_A)) + \mathcal{D}_A = \mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A + \mathcal{D}_A - \mathcal{P}'_A - \mathcal{D}_A + (\mathcal{P}_A \cap \mathcal{P}'_A) - \mathcal{D}_A + \mathcal{D}_A = (\mathcal{I}_A - \mathcal{P}_A + \mathcal{D}_A) \cap (\mathcal{I}_A - \mathcal{P}'_A + \mathcal{D}_A) = \overline{\mathcal{P}}_A \cap \overline{\mathcal{P}'_A}$. The second identity directly ensues. The inclusion relation $\subseteq$, defined by conditions (27), is by consequence reversed when the projectors it involves are negated. Indeed, if $\mathcal{P}'_A \subseteq \mathcal{P}_A \implies \mathcal{P}'_A \cap \mathcal{P}_A = \mathcal{P}'_A$, negating both sides of the cap yields $\overline{\mathcal{P}'_A \cap \mathcal{P}_A} = \overline{\mathcal{P}'_A} \cup \overline{\mathcal{P}_A}$. But the term under the negation in $\overline{\mathcal{P}'_A \cap \mathcal{P}_A}$ is equivalent to $\mathcal{P}_A$ since the inclusion also implies that $\mathcal{P}'_A \subseteq \mathcal{P}_A \implies \mathcal{P}'_A \cup \mathcal{P}_A = \mathcal{P}_A$. Whence, $\overline{\mathcal{P}'_A} \cap \overline{\mathcal{P}_A} = \overline{\mathcal{P}'_A} \cup \overline{\mathcal{P}_A} = \overline{\mathcal{P}_A}$, which by definition is the inclusion $\overline{\mathcal{P}_A} \subseteq \overline{\mathcal{P}'_A}$. This reasoning works in both ways, proving

$$\mathcal{P}'_A \subseteq \mathcal{P}_A \iff \overline{\mathcal{P}}_A \subseteq \overline{\mathcal{P}'_A}, \quad (120)$$

which is equation (31) in the main text.

A.4 Adding tensor and transformation almost makes Linear Logic

A.4.1 The tensor operation

For Definition 5, we also defined a tensor product at the level of the projectors. Here we now present its algebraic properties so that we can meaningfully consider $\mathcal{A} \otimes \mathcal{B} \subset \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ when we know that $\mathcal{A}$ is characterized by a projector $\mathcal{P}_A$ acting on $\mathcal{L}(\mathcal{H}^A)$ and that $\mathcal{B}$ is by $\mathcal{P}_B$. Because of the isomorphism $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B) \cong$

$\mathcal{L}(\mathcal{H}^A) \otimes \mathcal{L}(\mathcal{H}^B)$, the definition of **no signaling composition**, nicknamed **tensor**, is straightforward:

$$(\mathcal{P}_A \otimes \mathcal{P}_B) \left\{ \sum_i q_i (V_i \otimes U_i) \right\} \equiv \sum_i q_i (\mathcal{P}_A \{V_i\} \otimes \mathcal{P}_B \{U_i\}) , \quad (121)$$

and it should hold for all i such that $q_i \in \mathbb{C}$, $V_i \in \mathcal{L}(\mathcal{H}^A)$, $U_i \in \mathcal{L}(\mathcal{H}^B)$ as any operator in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ can be expressed as such.

By the inherited properties of the tensor product at the level of operators, the tensor at the level of superoperators is associative; it is moreover distributive with respect to the cap,

$$(\mathcal{P}_A \cap \mathcal{P}'_A) \otimes \mathcal{P}_B = (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}'_A \otimes \mathcal{P}_B) , \quad (122)$$

as well as to the cup,

$$(\mathcal{P}_A \cup \mathcal{P}'_A) \otimes \mathcal{P}_B = (\mathcal{P}_A \otimes \mathcal{P}_B) \cup (\mathcal{P}'_A \otimes \mathcal{P}_B) . \quad (123)$$

This directly follows from idempotency:

$$(\mathcal{P}_A \cap \mathcal{P}'_A) \otimes \mathcal{P}_B = (\mathcal{P}_A \cap \mathcal{P}'_A) \otimes (\mathcal{P}_B \cap \mathcal{P}_B) , \quad (124)$$

and because the intersection of projectors is party-wise, thus it commutes with the tensor product. Another way to see it is that composition of superoperators obeys an interchange law with the tensor product of superoperators:

$$\begin{aligned} & (\mathcal{P}_A \cap \mathcal{P}'_A) \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) \\ &= (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}'_A \otimes \mathcal{P}'_B) . \end{aligned} \quad (125)$$

Tensor preserves projectors on operator systems: Expressions built from the tensor product of Hermitian-preserving projectors are automatically HP since the dagger distributes over the tensor, $(\rho \otimes \sigma)^\dagger = \rho^\dagger \otimes \sigma^\dagger$. If the composed projectors are self-adjoint then so is their tensor product since the inner product is $\mathbb{C}$ -linear and splits as $\langle \mathcal{P}_A \{ \rho_A \} \otimes \mathcal{P}_B \{ \sigma_B \} , \eta_A \otimes \chi_B \rangle_{AB} = \langle \mathcal{P}_A \{ \rho_A \} , \eta_A \rangle_A \langle \mathcal{P}_B \{ \sigma_B \} , \chi_B \rangle_B$ so that $(\mathcal{P}_A \otimes \mathcal{P}_B)^* = \mathcal{P}_A^* \otimes \mathcal{P}_B^*$. They preserve idempotency,

$$\begin{aligned} & (\mathcal{P}_A \otimes \mathcal{P}_B)^2 = (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{P}_B) \\ & \stackrel{(125)}{=} (\mathcal{P}_A \cap \mathcal{P}_A) \otimes (\mathcal{P}_B \cap \mathcal{P}_B) , \end{aligned} \quad (126)$$

because of interchange law, thus

$$(\mathcal{P}_A \otimes \mathcal{P}_B)^2 = (\mathcal{P}_A \otimes \mathcal{P}_B) . \quad (127)$$

A further consequence of the interchange law is that the tensor composition of, respectively, the units and the zeroes on A and B , $\mathcal{I}_A \otimes \mathcal{I}_B$ and $\mathcal{D}_A \otimes \mathcal{D}_B$, respectively yield the unit and zero of the projector algebra on $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$, $\mathcal{I}_{AB}$ and $\mathcal{D}_{AB}$. Indeed, they can be used to define the negation of $\mathcal{P}_A \otimes \mathcal{P}_B$, $\overline{\mathcal{P}_A \otimes \mathcal{P}_B}$, and this gives the correct Boolean completions: $(\mathcal{P}_A \otimes \mathcal{P}_B) \cap \overline{(\mathcal{P}_A \otimes \mathcal{P}_B)} = (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B) = (\mathcal{P}_A \otimes \mathcal{P}_B - \mathcal{P}_A^2 \otimes \mathcal{P}_B^2 + \mathcal{D}_A \otimes \mathcal{D}_B) = \mathcal{D}_A \otimes \mathcal{D}_B$ (at the second equality sign, the interchange law was used to compute the cap of $\mathcal{P}_A \otimes \mathcal{P}_B$ with each of the three components of the negation). $(\mathcal{P}_A \otimes \mathcal{P}_B) \cup \overline{(\mathcal{P}_A \otimes \mathcal{P}_B)} = (\mathcal{P}_A \otimes \mathcal{P}_B) \cup (\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B) = \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \mathcal{D}_B = \mathcal{I}_A \otimes \mathcal{I}_B$. Hence,

$$(\mathcal{P}_A \otimes \mathcal{P}_B) \cap \overline{(\mathcal{P}_A \otimes \mathcal{P}_B)} = \mathcal{D}_A \otimes \mathcal{D}_B \equiv \mathcal{D}_{AB} ; \quad (128a)$$

$$(\mathcal{P}_A \otimes \mathcal{P}_B) \cup \overline{(\mathcal{P}_A \otimes \mathcal{P}_B)} = \mathcal{I}_A \otimes \mathcal{I}_B \equiv \mathcal{I}_{AB} . \quad (128b)$$

Tensor composition then “preserves the identity” simply because the identity on a joint system is the tensor product of the identities on each of the spaces being combined, $\mathbb{1}_{AB} = \mathbb{1}_A \otimes \mathbb{1}_B$. In again because of the interchange law, $\mathcal{D}_{AB} \cap (\mathcal{P}_A \otimes \mathcal{P}_B) = (\mathcal{D}_A \otimes \mathcal{D}_B) \cap (\mathcal{P}_A \otimes \mathcal{P}_B) = (\mathcal{D}_A \cap \mathcal{P}_A) \otimes (\mathcal{D}_B \cap \mathcal{P}_B) = \mathcal{D}_A \otimes \mathcal{D}_B$,

$$\mathcal{D}_{AB} \cap (\mathcal{P}_A \otimes \mathcal{P}_B) = \mathcal{D}_{AB} . \quad (129)$$

This shows that the tensor product of two projectors satisfying Eqs. (84) satisfies them as well.

Tensor product preserves commutation:

Since, if $\mathcal{P}_A \cap \mathcal{P}'_A = \mathcal{P}'_A \cap \mathcal{P}_A$ and $\mathcal{P}_B \cap \mathcal{P}'_B = \mathcal{P}'_B \cap \mathcal{P}_B$, then $(\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}'_A \otimes \mathcal{P}'_B) = (\mathcal{P}_A \cap \mathcal{P}'_A) \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) = (\mathcal{P}'_A \cap \mathcal{P}_A) \otimes (\mathcal{P}'_B \cap \mathcal{P}_B)$, and so

$$\begin{aligned} & (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}'_A \otimes \mathcal{P}'_B) \\ &= (\mathcal{P}'_A \otimes \mathcal{P}'_B) \cap (\mathcal{P}_A \otimes \mathcal{P}_B) . \end{aligned} \quad (130)$$

Negation and tensor: Interestingly, however, the tensor and negation do not commute with each other: $\overline{\mathcal{P}_A \otimes \mathcal{P}_B} \neq \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}$. Actually, the latter expression defines a subspace in the former since

$$\begin{aligned} & \overline{\mathcal{P}_A \otimes \mathcal{P}_B} \cap \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \\ &= (\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \cap \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \\ &= \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} - \mathcal{D}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B, \quad (131) \end{aligned}$$

so that

$$\overline{\mathcal{P}_A \otimes \mathcal{P}_B} \cap \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} = \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}. \quad (132)$$

And from this result

$$\overline{\mathcal{P}_A \otimes \mathcal{P}_B} \cup \overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} = \overline{\mathcal{P}_A \otimes \mathcal{P}_B}. \quad (133)$$

Hence,

$$\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \subseteq \overline{\mathcal{P}_A \otimes \mathcal{P}_B}. \quad (134)$$

(Recall that the subset symbol, $\subseteq$, implicitly means that we are considering the images of these projectors.) This identity which, using (120), can be recast as

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}}, \quad (135)$$

is Eq. (33) in the main text. That these two projectors are not equal is fundamentally the reason why an indefinite causal order may arise; this will become clearer after the prec is introduced in the algebra so that equations (165) and (166) can be written down. For an illustration of this relation, the left-hand side of this identity is depicted in the leftmost diagram in Fig. 9, whereas the right-hand side is the rightmost.

Because of the interchange law (125), relation (134) is stable when composed with more parties. More generally, the inclusion relation $\subseteq$ induced by conditions (28) is stable under the tensor; if $\mathcal{P}_A \subseteq \mathcal{P}'_A$, then the following hold

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}'_A \otimes \mathcal{P}_B, \quad (136a)$$

$$\overline{\mathcal{P}_A \otimes \mathcal{P}_B} \subseteq \overline{\mathcal{P}'_A \otimes \mathcal{P}_B}, \quad (136b)$$

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \overline{\overline{\mathcal{P}'_A} \otimes \overline{\mathcal{P}_B}}. \quad (136c)$$

The last equation is a consequence of the first using (134). The first equation holds because

$$\begin{aligned} & (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}'_A \otimes \mathcal{P}_B) \\ & \stackrel{(122)}{=} (\mathcal{P}_A \cap \mathcal{P}'_A) \otimes \mathcal{P}_B, \quad (137) \end{aligned}$$

and the term in parenthesis is equal to $\mathcal{P}_A$ because $\mathcal{P}_A \subseteq \mathcal{P}'_A$, therefore the first equation holds. The second equation is proven using De Morgan duality:

$$\begin{aligned} & \overline{\mathcal{P}_A \otimes \mathcal{P}_B} \cap \overline{\mathcal{P}'_A \otimes \mathcal{P}_B} \\ & \stackrel{(119a)}{=} \overline{(\overline{\mathcal{P}_A \otimes \mathcal{P}_B}) \cup (\overline{\mathcal{P}'_A \otimes \mathcal{P}_B})} \\ & \stackrel{(123)}{=} \overline{(\overline{\mathcal{P}'_A \cup \mathcal{P}_A}) \otimes \overline{\mathcal{P}_B}} \\ & \stackrel{(119b)}{=} \overline{\mathcal{P}'_A \cap \mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \\ & = \overline{\mathcal{P}'_A} \otimes \overline{\mathcal{P}_B}. \quad (138) \end{aligned}$$

Relations (134) and (136) are important to define the no signaling subset of an arbitrary projector, as explained in section 4.2.2.

Finally, observe that the only time the identity $\mathcal{P}_A \otimes \mathcal{P}_B = \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}}$ holds is when the projector is either characterizing the totality of the space or only the span of identity. In equation,

$$\mathcal{P}_A \otimes \mathcal{P}_B = \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}} \iff \mathcal{P}_A = \mathcal{P}_B = \mathcal{I} \text{ or } \mathcal{D}. \quad (139)$$

This is proven by rewriting $\overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}}$ into

$$\begin{aligned} & \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}} = \mathcal{P}_A \otimes \mathcal{P}_B + (\overline{\mathcal{P}_A} \otimes \mathcal{P}_B \\ & \quad - \overline{\mathcal{P}_A} \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\ & \quad + (\mathcal{P}_A \otimes \overline{\mathcal{P}_B} \\ & \quad - \mathcal{P}_A \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \overline{\mathcal{P}_B} + \mathcal{D}_A \otimes \mathcal{D}_B), \quad (140) \end{aligned}$$

using the definition of negation and algebraic properties. It can then be understood from Fig. 9: the first term of the above is the right quarter of the wheel, the second is the top quarter with its boundary removed, and the third is the bottom also without boundaries. As the three parts share no intersection, the regular addition '+' is equivalent to a conjunction '∩'. Next, more algebraic manipulations lead to

$$\begin{aligned} & \overline{\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B}} = \mathcal{P}_A \otimes \mathcal{P}_B + (\overline{\mathcal{P}_A} \otimes \mathcal{P}_B \\ & \quad + \mathcal{P}_A \otimes \overline{\mathcal{P}_B} - \mathcal{I}_A \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \mathcal{I}_B), \quad (141) \end{aligned}$$

and from this expression it is direct to check that the term in parentheses vanishes if and only if either of the conditions in Eq. (139) holds.

A.4.2 The transformation operation

The projector appearing in Proposition 2 is built from the projectors characterizing its input and output. The corresponding operation is defined as the **transformation** in Eq. (57), represented by $\rightarrow$:

$$\begin{aligned} \mathcal{P}_A \rightarrow \mathcal{P}_B &\equiv \mathcal{I}_A \otimes \mathcal{I}_B \\ &- \mathcal{P}_A \otimes \mathcal{I}_B + \mathcal{P}_A \otimes \mathcal{P}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \end{aligned} \quad (142)$$

This additional operation in the Boolean algebra of projectors is actually secondary since it can be entirely defined using the negation and the no signaling composition (i.e., the tensor):

$$\mathcal{P}_A \rightarrow \mathcal{P}_B \equiv \overline{\mathcal{P}_A \otimes \overline{\mathcal{P}_B}}. \quad (143)$$

Therefore it will automatically be a valid projector in the sense that it will obey Eqs. (84) if its constituents do, and will accordingly preserve commutation if they do. Furthermore, as this expression involves both composition and negation, it will not be associative in general. For example, $((\mathcal{P}_A \rightarrow \mathcal{P}_B) \rightarrow \mathcal{P}_C) \neq (\mathcal{P}_A \rightarrow (\mathcal{P}_B \rightarrow \mathcal{P}_C))$, because the *uncurrying rule* proven in Ref. [4] can be applied to the right-hand side:

$$\mathcal{P}_A \rightarrow (\mathcal{P}_B \rightarrow \mathcal{P}_C) = (\mathcal{P}_A \otimes \mathcal{P}_B) \rightarrow \mathcal{P}_C, \quad (144)$$

which is obviously different than $((\mathcal{P}_A \rightarrow \mathcal{P}_B) \rightarrow \mathcal{P}_C)$ by Eq. (135). In the algebra, the uncurrying rule relies on the associativity of the tensor:

$$\begin{aligned} \mathcal{P}_A \rightarrow (\mathcal{P}_B \rightarrow \mathcal{P}_C) &= \overline{\mathcal{P}_A \otimes (\overline{\mathcal{P}_B \otimes \mathcal{P}_C})} \\ &= \overline{\mathcal{P}_A \otimes \overline{\overline{\mathcal{P}_B \otimes \mathcal{P}_C}}} \stackrel{(115)}{=} \overline{\mathcal{P}_A \otimes \mathcal{P}_B \otimes \mathcal{P}_C} \\ &= \overline{(\mathcal{P}_A \otimes \mathcal{P}_B) \otimes \mathcal{P}_C} = (\mathcal{P}_A \otimes \mathcal{P}_B) \rightarrow \mathcal{P}_C. \end{aligned} \quad (145)$$

Remark that because of Eq. (136b), the transformation also preserves the inclusion relations:

$$\mathcal{P}_A \subseteq \mathcal{P}'_A \Rightarrow (\mathcal{P}_A \rightarrow \mathcal{P}_B) \subseteq (\mathcal{P}'_A \rightarrow \mathcal{P}_B). \quad (146)$$

At the level of Boolean logic, the transformation operation can be understood as a logical implication. Indeed, the transformation is equal to its inverse implication, meaning that it satisfies

$$\mathcal{P}_A \rightarrow \mathcal{P}_B = \overline{\mathcal{P}_A} \leftarrow \overline{\mathcal{P}_B}, \quad (147)$$

which comes from the definition. Additionally, note that it is equivalent to $\overline{\mathcal{P}_B} \rightarrow \overline{\mathcal{P}_A}$ since the order of the systems in the tensor product does not matter as $\mathcal{H}^A \otimes \mathcal{H}^B \cong \mathcal{H}^B \otimes \mathcal{H}^A$, we only keep it fixed to avoid adding confusion.

The $\rightarrow$ operation on projectors recovers the $\rightarrow$ type constructor of Bisio and Perinotti type theory [4, 5] (see Sec. 2): if we define the trivial system as ‘1’, that is to say, the 1-dimensional state structure $\{1\}$ made of the number 1, one can interpret the measurement (thus the negation of a given state structure) as a transformation into the trivial system. This leads to the identity

$$\overline{\mathcal{P}_A} = \mathcal{P}_A \rightarrow 1, \quad (148)$$

which justifies the notation $\mathcal{A} \rightarrow 1 \equiv \overline{\mathcal{A}}$. The proof is straightforward from the definition since 1 is 1-dimensional: $\mathcal{P}_A \rightarrow 1 = \overline{\mathcal{P}_A \otimes \overline{1}} = \overline{\mathcal{P}_A}$. In the same way, one can prove that

$$\mathcal{P}_A = 1 \rightarrow \mathcal{P}_A. \quad (149)$$

Therefore, a state structure can be seen as a transformation from the trivial system to itself and its negation as a transformation from itself to the trivial system. In view of the link between composition and transformation, one may also interpret the former in terms of the latter. A bipartite system in tensor-composed state structures, $\mathcal{A} \otimes \mathcal{B}$ for instance, can actually be seen as characterized by the following transformations

$$\mathcal{P}_A \otimes \mathcal{P}_B = \overline{\mathcal{P}_A \rightarrow \overline{\mathcal{P}_B}} = \overline{\mathcal{P}_A} \leftarrow \overline{\mathcal{P}_B}. \quad (150)$$

This means that a no signaling composite bipartite system is equivalent to a functional on a transformation from one state structure to the functionals on the other of the other as explained in Sec. 2. In the above equation, we see that the direction of the transformation has no influence, which is expected since it is a no signaling composition; this interpretation is explained in Sec. 5.

A.4.3 Linear Logic

We now show under which definition that the rules $\{\overline{\cdot}, \cap, \cup, \otimes, \overline{\cdot} \rightarrow \cdot\}$ form a model of logic which is almost multiplicative additive linear logic (MALL). The correspondence with MALL is summarized in Table 2, where the rules we have introduced are put in correspondence with their usual notation.

Name	Symbol		Unit	
	proj.	LL	proj.	LL
Negation	$\overline{\cdot}$	$\cdot^\perp$	/	/
Additive conjunction	$\cap$	$\&$	$\mathcal{I}$	$\top$
Additive disjunction	$\cup$	$\oplus$	$\mathcal{D}$	0
Multiplicative conjunction	$\otimes$	$\otimes$	1	1
Multiplicative disjunction	$\overline{\cdot} \rightarrow \cdot$	$\wp$	1	$\perp$
Linear Implication	$\rightarrow$	$\multimap$	1	/

Table 2: Correspondence of the algebraic rules of the projective characterization (proj.) with linear logic (LL)

Linear logic is a formal system of logic, of which MALL is a fragment, that is a restriction of the logic to fewer rules. It can be defined as a *sequent calculus*, that is the formalization of proof systems in which each *proposition* follows under some *structural and inference rules* from other propositions (it is a *sequent* of these propositions, whence the name) see Ref. [14]. If one sees the projectors algebra and the operations on it as the propositions, the starting set of propositions is indeed similar to those of MALL:

1. The set from which we start is made of valid projectors;
2. For every projector $\mathcal{P}$, there exists a projector $\overline{\mathcal{P}}$;
3. For every projectors $\mathcal{P}_A$ and $\mathcal{P}'_A$, there exist an additive conjunction $\mathcal{P}_A \cap \mathcal{P}'_A$ as well as an additive disjunction $\mathcal{P}_A \cup \mathcal{P}'_A$;
4. For every projectors $\mathcal{P}_A$ and $\mathcal{P}_B$, there exist a multiplicative conjunction $\mathcal{P}_A \otimes \mathcal{P}_B$ as well as a multiplicative disjunction $\overline{\mathcal{P}}_A \rightarrow \mathcal{P}_B$;
5. There are two constants ($\mathcal{I}, \mathcal{D}$) that go with each additive binary connectors;
6. There are two constants ($1, 1$) that go with each multiplicative binary connectors.

And the properties that can be proven for the projectors are similar to those that can be proven in MALL using the sequent calculus:

1. All binary connectors are commutative;
2. Multiplicative connectors distribute over additive ones;
3. All propositions have a negation obeying the De Morgan rules (42);

4. Additive constants are the negation of each other;
5. Multiplicative constants are the negation of each other.

Proposition 1 is always true as we take it as an axiom. Propositions 2 to 4 are true from the fact that the definition of these various connectors implies closure, which, in the case of projectors, is the conservation of properties (84). These were proven for each connector in the previous sections. Proposition 5 follows from equations (106) and (109). Proposition 6 happens because of the isomorphism $\mathcal{L}(\mathcal{H}) \otimes \mathbb{C} \cong 1$ so that $\mathcal{P} \otimes 1 = \mathcal{P}$ and the same way, $\overline{\mathcal{P}} \otimes \overline{1} = \overline{\mathcal{P}} = \mathcal{P}$, $1 \otimes \overline{\mathcal{P}} = \overline{\mathcal{P}} = \mathcal{P}$.

Property 1 is true from the definition. In the case of the multiplicative connectors, $\mathcal{P}_A \otimes \mathcal{P}_B$ and $\overline{\mathcal{P}}_A \otimes \overline{\mathcal{P}}_B$, the isomorphism $\mathcal{H}^A \otimes \mathcal{H}^B \cong \mathcal{H}^B \otimes \mathcal{H}^A$ should of course be used.

Property 2 follows from Eqs. (122) and (123) in the case of $\otimes$ and application of the De Morgan rules (119) on these two equations can be used to prove the property in the case of $\overline{\cdot} \rightarrow \cdot$. Put differently, the multiplicative conjunction (respectively, disjunction) distributes over the additive conjunction (disjunction)

$$\begin{aligned} \mathcal{P}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) \\ = (\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{P}'_B) ; \end{aligned} \quad (151a)$$

$$\begin{aligned} \overline{\mathcal{P}}_A \rightarrow (\mathcal{P}_B \cup \mathcal{P}'_B) \\ = (\overline{\mathcal{P}}_A \rightarrow \mathcal{P}_B) \cup (\overline{\mathcal{P}}_A \rightarrow \mathcal{P}'_B) . \end{aligned} \quad (151b)$$

But as $\otimes$ and $\rightarrow$ are both operations that merge subspaces, the converse obviously does not hold. Indeed an expression like $\mathcal{P}_X \cap (\mathcal{P}_A \otimes \mathcal{P}_B) = (\mathcal{P}_X \otimes \mathcal{P}_A) \cap (\mathcal{P}_X \otimes \mathcal{P}_B)$ makes no sense as the

right-hand side feature a cap between two superoperator projectors defined on different spaces that are not isomorphic in general.

Property 3 was proven at Eq. (115) for the negation, Eq. (119) for the additive connectors, and follows from the definition in the case of multiplicative connectors.

Property 4 is the statement $\bar{\mathcal{I}} = \mathcal{I} - \mathcal{I} + \mathcal{D} = \mathcal{D}$, whose converse holds by idempotency or can be proven by the same kind of computation.

Property 5 is the statement $\bar{1} = 1 - 1 + 1$, as in the 1-dimensional case $\mathcal{D} = \mathcal{I} = 1$.

There are however some discrepancies with MALL, that we now discuss. First, two small issues that indicate a departure from a faithful model of MALL as in Ref. [14]:

1. The multiplicative units are equivalent;
2. The additive falsity (i.e. the unit for $\cup$) is not absorbing.

The first issue is why we use the terminology ‘degenerate’ in the main text to refer to the model of MALL formed by the algebra of projectors. The second issue –that $\mathcal{P}_A \otimes \mathcal{D}_B \neq \mathcal{D}_A \otimes \mathcal{D}_B$ – is more severe as it goes against an equality that can be proven in MALL. A way to circumvent this is to redefine the additive falsity as the number 0, but then the trace normalization of all operator systems the projectors characterize should be set to zero which would jeopardize the probabilistic interpretation.

A more substantial problem is the fact that the cap and cup are only well-defined for expressions featuring the same set of base projectors: two projectors can only be composed with a cap if they can be embedded in the same space. This is part of the problem we evoked when discussing Property 2 above. For the model of logic to work properly, the number of subsystems as well as the dimension of each subsystem must be fixed *a priori*, and each expression is associated with a given (set of) subsystem(s), so that the composition is well-defined. Knowing this information allows one to ‘pad’ the expressions with identities so as to embed them in the global Hilbert space. For example, an expression like $\mathcal{P}_X \cap (\mathcal{P}_A \otimes \mathcal{P}_B)$ can be made definite by knowing for example that there are two systems A and B , and that $\mathcal{P}_X$ is associated with system B so that the padding reads $\mathcal{P}_X \cong \mathcal{I}_A \otimes \mathcal{P}_B^X$, where $\mathcal{P}_B^X$ is the projector $\mathcal{P}_X$ acting on system B . Another possibility is

that $\mathcal{P}_X$ is associated with both systems, so that the padding is trivial $\mathcal{P}_X \cong \mathcal{P}_{AB}^X$. Yet another possibility is that the projector is associated with system A and B , but the global Hilbert space is tripartite: $\mathcal{H} = \mathcal{H}^A \otimes \mathcal{H}^B \otimes \mathcal{H}^C$. In that case, the padding should be done on the two projectors: $\mathcal{P}_X \cong \mathcal{P}_{AB}^X \otimes \mathcal{I}_C$, $\mathcal{P}_A \otimes \mathcal{P}_B \cong \mathcal{P}_A \otimes \mathcal{P}_B \otimes \mathcal{I}_C$.

Despite these particularities, the connection between the algebra of projectors and MALL is too obvious not to be pointed out.

Remark that the issues with interpreting the algebra as a model of LL come from the choice of additive connectives as $\cap$ and $\cup$. This choice is motivated because we want to compare the underlying operator systems associated with different types of higher-order quantum transformations. Another choice proposed in Ref. [26] is to use the Cartesian product and direct sum as the additive conjunction and disjunction respectively. In that case, the issues can all be alleviated except for the degeneracy of the multiplicative units... but the additive connectors can no longer be used for comparison.

A.5 Adding the prec

Define the **A-to-B one-way signaling composition**, nicknamed **prec**, as

$$\mathcal{P}_A \prec \mathcal{P}_B \equiv \mathcal{I}_A \otimes \mathcal{P}_B - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (152)$$

The reversed sign $\succ$ (also nicknamed the **prec**) can be defined accordingly: the **B-to-A one-way signaling composition** is given by

$$\mathcal{P}_A \succ \mathcal{P}_B \equiv \mathcal{P}_A \otimes \mathcal{I}_B - \mathcal{D}_A \otimes \bar{\mathcal{P}}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (153)$$

As is the case with the transformation $\rightarrow$, note that $\mathcal{P}_A \succ \mathcal{P}_B \cong \mathcal{P}_B \prec \mathcal{P}_A$ since $\mathcal{H}^A \otimes \mathcal{H}^B \cong \mathcal{H}^B \otimes \mathcal{H}^A$.

Prec preserves projectors on operator systems: When introduced in the main text, Eq. (50), this projector was obtained as the intersection of the following two projectors:

$$\mathcal{P}_A \prec \mathcal{P}_B \equiv (\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B). \quad (154)$$

Because $\mathcal{I}_A, \mathcal{P}_A, \mathcal{P}_B$ are projectors on operator systems, meaning they obey Eqs. (84), and we proved that operations $\bar{\cdot}, \otimes$ and $\cap$ preserve this property in the last sections, $\mathcal{P}_A \prec \mathcal{P}_B$ obeys Eqs. (84) as well and it is therefore a projector on an operator system in $\mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$.

Prec preserves commutation: For the same reason as above, if $\mathcal{P}_A \cap \mathcal{P}'_A = \mathcal{P}'_A \cap \mathcal{P}_A$ and $\mathcal{P}_B \cap \mathcal{P}'_B = \mathcal{P}'_B \cap \mathcal{P}_B$, then

$$\begin{aligned} (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}'_A \prec \mathcal{P}'_B) \\ = (\mathcal{P}'_A \prec \mathcal{P}'_B) \cap (\mathcal{P}_A \prec \mathcal{P}_B) . \end{aligned} \quad (155)$$

It also commutes with the B-to-A one-way signaling composition:

$$\begin{aligned} (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}'_A \succ \mathcal{P}'_B) \\ = (\mathcal{P}'_A \succ \mathcal{P}'_B) \cap (\mathcal{P}_A \prec \mathcal{P}_B) , \end{aligned} \quad (156)$$

again because

$$\mathcal{P}_A \succ \mathcal{P}_B \equiv \mathcal{P}_A \otimes \mathcal{I}_B - \mathcal{D}_A \otimes \bar{\mathcal{P}}_B + \mathcal{D}_A \otimes \mathcal{D}_B , \quad (157)$$

is an expression derived from connectives that preserve commutation:

$$\mathcal{P}_A \succ \mathcal{P}_B \equiv (\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{I}_B) . \quad (158)$$

Proving (156) thus reduces to prove the commutation of the terms in parenthesis in $((\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B)) \cap ((\mathcal{P}'_A \rightarrow \mathcal{P}'_B) \cap (\mathcal{P}'_A \otimes \mathcal{I}_B))$, which follows from the associativity of the $\cap$ and the commutation of the $\rightarrow$ with the $\otimes$. The same way,

$$\begin{aligned} (\mathcal{P}_A \succ \mathcal{P}_B) \cap (\mathcal{P}'_A \succ \mathcal{P}'_B) \\ = (\mathcal{P}'_A \succ \mathcal{P}'_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B) . \end{aligned} \quad (159)$$

A.5.1 The prec is associative and commutes with the negation

Compared to transformation, one-way signaling composition is better behaved in the sense that it is associative

$$\mathcal{P}_A \prec (\mathcal{P}_B \prec \mathcal{P}_C) = (\mathcal{P}_A \prec \mathcal{P}_B) \prec \mathcal{P}_C . \quad (160)$$

So that $\mathcal{P}_A \prec (\mathcal{P}_B \prec \mathcal{P}_C) = \mathcal{P}_A \prec \mathcal{P}_B \prec \mathcal{P}_C$ can be written unambiguously. Even more, it is better behaved than both the no signaling (the tensor) and the two-way signaling (the transformation) compositions as it is distributive with respect to the negation

$$\overline{\mathcal{P}_A \prec \mathcal{P}_B} = \bar{\mathcal{P}}_A \prec \bar{\mathcal{P}}_B . \quad (161)$$

We first prove the distributivity of the negation by direct computation,

$$\begin{aligned} \overline{\mathcal{P}_A \prec \mathcal{P}_B} \\ = \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \prec \mathcal{P}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\ = \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{I}_A \otimes \mathcal{P}_B + \bar{\mathcal{P}}_A \otimes \mathcal{D}_B \\ \quad - \mathcal{D}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\ = (\mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{I}_A \otimes \mathcal{P}_B + \mathcal{I}_A \otimes \mathcal{D}_B) \\ \quad - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\ = (\mathcal{I}_A \otimes \bar{\mathcal{P}}_B) - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\ = \bar{\mathcal{P}}_A \prec \bar{\mathcal{P}}_B . \end{aligned} \quad (162)$$

This can be used to prove associativity,

$$\begin{aligned} \mathcal{P}_A \prec (\mathcal{P}_B \prec \mathcal{P}_C) \\ = \mathcal{I}_A \otimes (\mathcal{P}_B \prec \mathcal{P}_C) \\ \quad - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C + \mathcal{D}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C \\ = \mathcal{I}_A \otimes \mathcal{I}_B \otimes \mathcal{P}_C - \mathcal{I}_A \otimes \bar{\mathcal{P}}_B \otimes \mathcal{D}_C + \mathcal{I}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C \\ \quad - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C + \mathcal{D}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C \\ = (\mathcal{I}_A \otimes \mathcal{I}_B) \otimes \mathcal{P}_C + \mathcal{D}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C \\ \quad - (\mathcal{I}_A \otimes \bar{\mathcal{P}}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \otimes \mathcal{D}_C \\ = (\mathcal{I}_A \otimes \mathcal{I}_B) \otimes \mathcal{P}_C \\ \quad - (\bar{\mathcal{P}}_A \prec \bar{\mathcal{P}}_B) \otimes \mathcal{D}_C + \mathcal{D}_A \otimes \mathcal{D}_B \otimes \mathcal{D}_C \\ = (\mathcal{I}_A \otimes \mathcal{I}_B) \otimes \mathcal{P}_C \\ \quad - \overline{(\mathcal{P}_A \prec \mathcal{P}_B)} \otimes \mathcal{D}_C + (\mathcal{D}_A \otimes \mathcal{D}_B) \otimes \mathcal{D}_C \\ = (\mathcal{P}_A \prec \mathcal{P}_B) \prec \mathcal{P}_C , \end{aligned} \quad (163)$$

where we have used the definition to go to the last line, and commutation of the prec with the negation to go to the penultimate one.

A.5.2 Inclusion relations of the prec

In Sec. A.4.1, it was shown that the tensor product of two projectors defines a subset of the composition of the same two projectors using the transformation connective seen as a composition, proving Eq. (33) of the main text. That is, $\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B = \overline{\bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B}$, see the discussion around Eq. (135). As the prec is itself a form of composition as argued in Sec. 5.1, similar inclusion relations can be proven. Like the tensor product, one-way signaling composition with the trivial system amounts to doing nothing since $\mathcal{P}_A \prec 1 = \mathcal{I}_A \otimes 1 - \bar{\mathcal{P}}_A \otimes 1 + \mathcal{D}_A \otimes 1 = \bar{\mathcal{P}}_A \otimes 1$ and $1 \prec \mathcal{P}_A = \mathcal{P}_A \otimes 1 - \mathcal{D}_A \otimes 1 + \mathcal{D}_A \otimes 1$ so that

$$\mathcal{P}_A \prec 1 = \mathcal{P}_A , \quad (164a)$$

$$1 \prec \mathcal{P}_A = \mathcal{P}_A . \quad (164b)$$

One can link the two-way signaling, one-way signaling, and no signaling compositions, respectively, $\bar{\cdot} \rightarrow \cdot$, $\cdot \prec \cdot$, and $\cdot \otimes \cdot$, by noticing that

$$(\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B) = \mathcal{P}_A \otimes \mathcal{P}_B, \quad (165)$$

and

$$(\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}_A \succ \mathcal{P}_B) = \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B. \quad (166)$$

Negating $\mathcal{P}_A$ yields relations (56) in the main text. Notice that the input of the transformation has to be negated to be interpreted as a composition because of how it was defined in Ref. [4]: it transforms an input in $\mathcal{L}(\mathcal{H}^A)$ to an output in $\mathcal{L}(\mathcal{H}^B)$, hence it must be a functional on $\mathcal{L}(\mathcal{H}^A)$, meaning it should locally belong to $\bar{\mathcal{A}}$ instead of $\mathcal{A}$ (see the discussion in Sec. 4.2.2). On the contrary, the tensor and the prec compose an output in $\mathcal{L}(\mathcal{H}^A)$ with an output in $\mathcal{L}(\mathcal{H}^B)$.

The first equation is proven by developing it, $(\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B) = (\mathcal{I}_A \otimes \mathcal{P}_B - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \cap (\mathcal{P}_A \otimes \mathcal{I}_B - \mathcal{D}_A \otimes \bar{\mathcal{P}}_B + \mathcal{D}_A \otimes \mathcal{D}_B)$, and noting that the intersection of any two elements but $(\mathcal{I}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{I}_B) = \mathcal{P}_A \otimes \mathcal{P}_B$ is giving $\mathcal{D}_A \otimes \mathcal{D}_B$. Thus, the expression reduces to $\mathcal{P}_A \otimes \mathcal{P}_B$ followed by eight occurrences of $\mathcal{D}_A \otimes \mathcal{D}_B$ alternating between a plus and minus sign, therefore canceling each other. The second equation has a quick proof using the De Morgan rule and the distributivity of the negation:

$$\begin{aligned} & (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}_A \succ \mathcal{P}_B) \\ & \stackrel{(115)}{=} \overline{(\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B)} \\ & \stackrel{(119a)}{=} \overline{(\mathcal{P}_A \prec \mathcal{P}_B) \cap (\bar{\mathcal{P}}_A \succ \bar{\mathcal{P}}_B)} \\ & \stackrel{(161)}{=} \overline{(\bar{\mathcal{P}}_A \prec \bar{\mathcal{P}}_B) \cap (\bar{\mathcal{P}}_A \succ \bar{\mathcal{P}}_B)} \\ & \stackrel{(165)}{=} \bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B = \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B. \end{aligned} \quad (167)$$

From Eq. (165), it can also be inferred that

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \prec \mathcal{P}_B, \quad (168a)$$

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \succ \mathcal{P}_B, \quad (168b)$$

as, obviously, $(\mathcal{P}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \prec \mathcal{P}_B) = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B) \cap (\mathcal{P}_A \prec \mathcal{P}_B) = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \succ \mathcal{P}_B) = \mathcal{P}_A \otimes \mathcal{P}_B$, and the second equation is proven analogously. The same way, from Eq. (166),

$$\mathcal{P}_A \prec \mathcal{P}_B \subseteq \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B, \quad (169a)$$

$$\mathcal{P}_A \succ \mathcal{P}_B \subseteq \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B. \quad (169b)$$

Putting these relations together,

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \prec \mathcal{P}_B \subseteq \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B, \quad (170a)$$

$$\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \succ \mathcal{P}_B \subseteq \bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B. \quad (170b)$$

The first line is diagrammatically depicted in

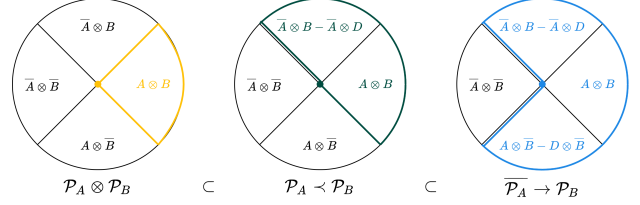


Figure 9: Diagram of the subspaces defined by the three ways of defining a composite projector.

Fig. 9: the subspace associated with $\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B$ is represented in blue on the right. The intermediate case, associated with $\mathcal{P}_A \prec \mathcal{P}_B$ is depicted in green on the centre; from Eq. (50), the green part, $\mathcal{P}_A \prec \mathcal{P}_B \equiv (\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B) \cap (\mathcal{I}_A \otimes \mathcal{P}_B)$, is obtained as the intersection of the blue part, $\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B$, with $\mathcal{I}_A \otimes \mathcal{P}_B = (\mathcal{P}_A \otimes \mathcal{P}_B) \cup (\bar{\mathcal{P}}_A \otimes \mathcal{P}_B)$ which correspond to the right and top quadrants. It indeed recovers the other definition (152), $\mathcal{P}_A \prec \mathcal{P}_B \equiv \mathcal{I}_A \otimes \mathcal{P}_B - \bar{\mathcal{P}}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B$: the green part made of the right and top quadrants $(\mathcal{I}_A \otimes \mathcal{P}_B)$ minus the line separating the left and top quadrants $(\bar{\mathcal{P}}_A \otimes \mathcal{D}_B)$ except at the central dot $(\mathcal{D}_A \otimes \mathcal{D}_B)$. The smallest case, associated with $\mathcal{P}_A \prec \mathcal{P}_B$ is depicted in yellow on the left; as discussed in Corollary 2 it is effectively defined by the intersection of no signaling subspaces of Lemma 5, $(\mathcal{I}_A \otimes \mathcal{P}_B) \cap (\mathcal{P}_A \otimes \mathcal{I}_B)$, that is the intersection of the right and top quadrants $(\mathcal{I}_A \otimes \mathcal{P}_B)$ with the right and bottom ones $(\mathcal{P}_A \otimes \mathcal{I}_B)$; the intersection with $\bar{\mathcal{P}}_A \rightarrow \mathcal{P}_B$ brings no new constraints.

Note that similar diagrams appear in Figs. 7 and 10, but they once again correspond to the situation where $\mathcal{P}_A$ is negated. Negating $\mathcal{P}_A$ in Eqs. (170) indeed yields the relations (60) of the main text depicted in these figures.

A.5.3 Distribution properties of the prec

The cap and prec satisfy an interchange law,

$$\begin{aligned} & (\mathcal{P}_A \cap \mathcal{P}'_A) \prec (\mathcal{P}_B \cap \mathcal{P}'_B) \\ & = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}'_A \prec \mathcal{P}'_B). \end{aligned} \quad (171)$$

This can be shown as follows:

$$\begin{aligned}
& (\mathcal{P}_A \cap \mathcal{P}'_A) \prec (\mathcal{P}_B \cap \mathcal{P}'_B) \\
&= \mathcal{I}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) - (\overline{\mathcal{P}_A \cap \mathcal{P}'_A}) \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\
&\stackrel{(119b)}{=} \mathcal{I}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) - (\overline{\mathcal{P}_A} \cup \overline{\mathcal{P}'_A}) \otimes \mathcal{D}_B \\
&\quad + \mathcal{D}_A \otimes \mathcal{D}_B \\
&= \mathcal{I}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) - \overline{\mathcal{P}_A} \otimes \mathcal{D}_B - \overline{\mathcal{P}'_A} \otimes \mathcal{D}_B \\
&\quad + (\overline{\mathcal{P}_A} \cap \overline{\mathcal{P}'_A}) \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\
&= (\mathcal{I}_A \otimes \mathcal{P}_B - \overline{\mathcal{P}_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&\quad \cap (\mathcal{I}_A \otimes \mathcal{P}'_B - \overline{\mathcal{P}'_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&= (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}'_A \prec \mathcal{P}'_B). \quad (172)
\end{aligned}$$

Remark that the grouping at the penultimate line is arbitrary. In different terms, because the cap is commutative we have $(\mathcal{P}_A \cap \mathcal{P}'_A) \prec (\mathcal{P}_B \cap \mathcal{P}'_B) = (\mathcal{P}_A \cap \mathcal{P}'_A) \prec (\mathcal{P}'_B \cap \mathcal{P}_B) = (\mathcal{P}_A \prec \mathcal{P}'_B) \cap (\mathcal{P}'_A \prec \mathcal{P}_B)$. Hence, the exact grouping does not matter as long as the A 's and B 's are on the correct side of the prec connector.

The cup and the prec satisfy an interchange law as well,

$$\begin{aligned}
& (\mathcal{P}_A \cup \mathcal{P}'_A) \prec (\mathcal{P}_B \cup \mathcal{P}'_B) \\
&= (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}'_A \prec \mathcal{P}'_B). \quad (173)
\end{aligned}$$

This can be shown the same way,

$$\begin{aligned}
& (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}'_A \prec \mathcal{P}'_B) \\
&= (\mathcal{I}_A \otimes \mathcal{P}_B - \overline{\mathcal{P}_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&\quad \cup (\mathcal{I}_A \otimes \mathcal{P}'_B - \overline{\mathcal{P}'_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&= (\mathcal{I}_A \otimes \mathcal{P}_B - \overline{\mathcal{P}_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&\quad + (\mathcal{I}_A \otimes \mathcal{P}'_B - \overline{\mathcal{P}'_A} \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B) \\
&\quad - \mathcal{I}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) + \overline{\mathcal{P}_A} \otimes \mathcal{D}_B + \overline{\mathcal{P}'_A} \otimes \mathcal{D}_B \\
&\quad - (\overline{\mathcal{P}_A} \cap \overline{\mathcal{P}'_A}) \otimes \mathcal{D}_B - \mathcal{D}_A \otimes \mathcal{D}_B \\
&= \mathcal{I}_A \otimes \mathcal{P}_B + \mathcal{I}_A \otimes \mathcal{P}'_B - \mathcal{I}_A \otimes (\mathcal{P}_B \cap \mathcal{P}'_B) \\
&\quad - (\overline{\mathcal{P}_A} \cap \overline{\mathcal{P}'_A}) \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\
&= \mathcal{I}_A \otimes (\mathcal{P}_B \cup \mathcal{P}'_B) - (\overline{\mathcal{P}_A \cup \mathcal{P}'_A}) \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B \\
&= (\mathcal{P}_A \cup \mathcal{P}'_A) \prec (\mathcal{P}_B \cup \mathcal{P}'_B). \quad (174)
\end{aligned}$$

These two interchange laws imply distributivity as a special case because the elements in the algebra are all idempotent, so that $\mathcal{P} = \mathcal{P} \cap \mathcal{P} = \mathcal{P} \cup \mathcal{P}$ and we can do the following: $(\mathcal{P}_A \cup \mathcal{P}'_A) \prec \mathcal{P}_B = (\mathcal{P}_A \cup \mathcal{P}'_A) \prec (\mathcal{P}_B \cup \mathcal{P}_B) =$

$(\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}'_A \prec \mathcal{P}_B)$. Therefore, the following relations hold:

$$(\mathcal{P}_A \cap \mathcal{P}'_A) \prec \mathcal{P}_B = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}'_A \prec \mathcal{P}_B), \quad (175a)$$

$$(\mathcal{P}_A \cup \mathcal{P}'_A) \prec \mathcal{P}_B = (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}'_A \prec \mathcal{P}_B), \quad (175b)$$

$$\mathcal{P}_A \prec (\mathcal{P}_B \cup \mathcal{P}'_B) = (\mathcal{P}_A \prec \mathcal{P}_B) \cup (\mathcal{P}_A \prec \mathcal{P}'_B), \quad (175c)$$

$$\mathcal{P}_A \prec (\mathcal{P}_B \cap \mathcal{P}'_B) = (\mathcal{P}_A \prec \mathcal{P}_B) \cap (\mathcal{P}_A \prec \mathcal{P}'_B). \quad (175d)$$

B Accidental isomorphisms in the case of quantum theory

The set inclusions (170), appearing in the main text as Eqs. (60):

$$\overline{\mathcal{P}_A} \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}_A} \prec \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B, \quad (176a)$$

$$\overline{\mathcal{P}_A} \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}_A} \succ \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B, \quad (176b)$$

can become equivalences in the special cases where the projectors are either identity or depolarizing. These imply set isomorphisms that are not without consequences: subsets with different signaling constraints get accidentally equivalent.

Putting an identity on the right side of the $\rightarrow$ gives

$$\mathcal{P}_A \rightarrow \mathcal{I}_B = \overline{\mathcal{P}_A} \otimes \mathcal{D}_B, \quad (177)$$

which happens to be equivalent to

$$\overline{\mathcal{P}_A} \prec \mathcal{I}_B = \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{P}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (178)$$

As it can be shown directly from the definitions. In that case, the two-way and one-way signaling compositions coincide. The same way, putting one on the left side gives

$$\begin{aligned}
\mathcal{I}_A \rightarrow \mathcal{P}_B &= \overline{\mathcal{I}_A} \otimes \overline{\mathcal{P}_B} \\
&= \mathcal{I}_A \otimes \mathcal{I}_B - \mathcal{I}_A \otimes \overline{\mathcal{P}_B} + \mathcal{D}_A \otimes \mathcal{D}_B \\
&= \mathcal{I}_A \otimes \mathcal{P}_B - \mathcal{I}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B, \quad (179)
\end{aligned}$$

which is equivalent to

$$\mathcal{D}_A \prec \mathcal{P}_B = \mathcal{I}_A \otimes \mathcal{P}_B - \mathcal{I}_A \otimes \mathcal{D}_B + \mathcal{D}_A \otimes \mathcal{D}_B. \quad (180)$$

One has then the following identities:

$$\overline{\mathcal{P}_A} \prec \mathcal{I}_B = \mathcal{P}_A \rightarrow \mathcal{I}_B; \quad (181a)$$

$$\mathcal{D}_A \prec \mathcal{P}_B = \mathcal{I}_A \rightarrow \mathcal{P}_B. \quad (181b)$$

These are the only two cases for which the transformation is equivalent to the prec as the following rewriting shows:

$$\mathcal{P}_A \rightarrow \mathcal{P}_B = \overline{\mathcal{P}}_A \prec \mathcal{P}_B + (\mathcal{I}_A - \mathcal{P}_A) \otimes (\mathcal{I}_B - \mathcal{P}_B). \quad (182)$$

These relations can be concisely recast as

$$\begin{aligned} \overline{\mathcal{P}}_A \prec \mathcal{P}_B &= \mathcal{P}_A \rightarrow \mathcal{P}_B \\ \iff \mathcal{P}_A &= \mathcal{I}_A \text{ or } \mathcal{P}_B = \mathcal{I}_B. \end{aligned} \quad (183)$$

In addition to that, one-way signaling composition is equivalent to the no signaling composition in the following cases:

$$\mathcal{P}_A \prec \mathcal{D}_B = \mathcal{P}_A \otimes \mathcal{D}_B; \quad (184a)$$

$$\mathcal{I}_A \prec \mathcal{P}_B = \mathcal{I}_A \otimes \mathcal{P}_B. \quad (184b)$$

Again, this directly follows from the definition by re-expressing it as

$$\overline{\mathcal{P}}_A \prec \mathcal{P}_B = \overline{\mathcal{P}}_A \otimes \mathcal{P}_B + (\mathcal{P}_A - \mathcal{D}_A) \otimes (\mathcal{P}_B - \mathcal{D}_B), \quad (185)$$

and this can be concisely recast as

$$\begin{aligned} \overline{\mathcal{P}}_A \prec \mathcal{P}_B &= \overline{\mathcal{P}}_A \otimes \mathcal{P}_B \\ \iff \mathcal{P}_A &= \mathcal{D}_A \text{ or } \mathcal{P}_B = \mathcal{D}_B. \end{aligned} \quad (186)$$

It should be noted that condition (186) is stronger than (183). Actually, when both conditions are satisfied at once, one has either of the two identities:

$$\mathcal{I}_A \rightarrow \mathcal{D}_B = \mathcal{D}_A \prec \mathcal{D}_B = \mathcal{D}_A \otimes \mathcal{D}_B; \quad (187a)$$

$$\mathcal{D}_A \rightarrow \mathcal{I}_A = \mathcal{I}_A \prec \mathcal{I}_B = \mathcal{I}_A \otimes \mathcal{I}_B. \quad (187b)$$

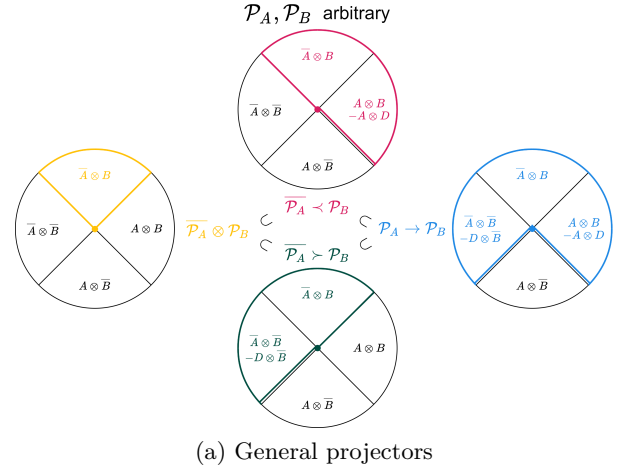
The reason this is the case comes from isomorphism (139), which reduces the transformation into a no signaling composition. Indeed,

$$\begin{aligned} \mathcal{D}_A \prec \mathcal{D}_B &\stackrel{(183)}{=} \mathcal{I}_A \rightarrow \mathcal{D}_B \stackrel{(143)}{=} \overline{\mathcal{I}_A \otimes \mathcal{D}_B} \\ &= \overline{\mathcal{I}_A \otimes \mathcal{I}_B} \stackrel{(139)}{=} \mathcal{D}_A \otimes \mathcal{D}_B. \end{aligned} \quad (188)$$

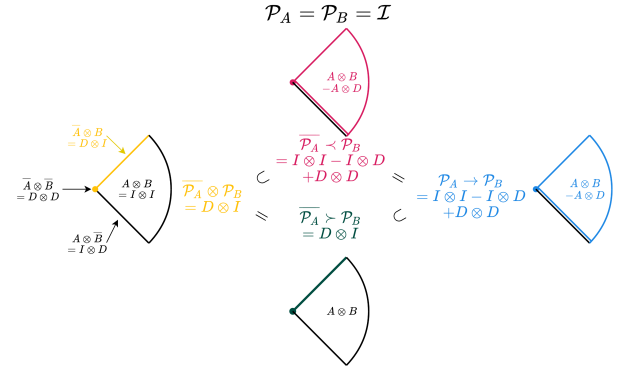
And the same way,

$$\begin{aligned} \mathcal{I}_A \prec \mathcal{I}_B &\stackrel{(183)}{=} \mathcal{D}_A \rightarrow \mathcal{I}_B \stackrel{(143)}{=} \overline{\mathcal{D}_A \otimes \mathcal{I}_B} \\ &= \overline{\mathcal{D}_A \otimes \mathcal{D}_B} \stackrel{(139)}{=} \mathcal{I}_A \otimes \mathcal{I}_B. \end{aligned} \quad (189)$$

Therefore, the isomorphisms (139), (183), and (186) give the conditions for set equivalences in the composition rules.



(a) General projectors



(b) Quantum theory (identity projectors)

Figure 10: Diagrams depicting the subspaces spanned by the four different ways of defining a transformation (top). When $\overline{\mathcal{P}}_A = \mathcal{D}_A$ and $\mathcal{P}_B = \mathcal{I}_B$ (bottom), the yellow and green zones are shrunk to the segment $\mathcal{D} \otimes B$: Equations akin to (184a) and (184b) are simultaneously satisfied so no signaling (yellow) is equivalent to one-way to A (green). At the same time, the pink and blue zones are reduced to $A \otimes B - A \otimes D$: Eqs. (181a) and (181b) are simultaneously satisfied so two-way signaling (blue) is equivalent to one-way signaling to B (pink).

One can understand these relations using a diagram of the kind depicted in Fig. 10. For a transformation $\mathcal{P}_A \rightarrow \mathcal{P}_B$ made of arbitrary projectors (Fig. 10a, in blue), its different substructures are effectively distinct (no signaling in yellow, A -to- B one-way in pink, and B -to- A one-way in green). Recall that one can infer the inclusion relations from their overlap, e.g. $\overline{\mathcal{P}}_A \otimes \mathcal{P}_B \subseteq \overline{\mathcal{P}}_A \prec \mathcal{P}_B \subseteq \mathcal{P}_A \rightarrow \mathcal{P}_B$ can be inferred from the fact that the yellow part is contained within the pink one which is contained within the blue one. Another example, $(\overline{\mathcal{P}}_A \prec \mathcal{P}_B) \cap (\overline{\mathcal{P}}_A \succ \mathcal{P}_B) = \mathcal{P}_A \rightarrow \mathcal{P}_B$ can be inferred from the fact that zone covered by the pink and green areas is equivalent to the blue one.

An accidental isomorphism is present when either the input or the output is a density operator (featuring an identity projector). In the first case, $\mathcal{P}_A = \mathcal{I}_A$ and $\overline{\mathcal{P}}_A = \mathcal{D}_A$, so the area $\overline{A}$ shrinks into a line: the top and left quadrants of the circle are shrunk into the diagonal that goes from bottom left to top right. The yellow zone is thus shrunk into $D \otimes B$, the upper border of $A \otimes B$, and so is the green zone, this equivalence of yellow and green is similar to Eq. (184a), but with A and B swapped: $\overline{\mathcal{P}}_A = \mathcal{D}_A \Rightarrow \overline{\mathcal{P}}_A \succ \mathcal{P}_B = \mathcal{D}_A \otimes \mathcal{P}_B$. At the same time, the pink zone is shrunk to $A \otimes B - A \otimes D$, i.e. $A \otimes B$ without its bottom border, and so is the blue zone, this equivalence of pink and blue is Eq. (181b). In the second case, $\mathcal{P}_B = \mathcal{I}_B$ and $\overline{\mathcal{P}}_B = \mathcal{D}_B$, so the area $\overline{B}$ also shrinks to a line: the bottom and left quadrants are shrunk into the diagonal that goes from top left to bottom right. The yellow zone is untouched but the green one is shrunk into the same zone as $\overline{A} \otimes \overline{B}$ becomes the border $\overline{A} \otimes D$, i.e. the top-left segment. This equivalence when $\mathcal{P}_B = \mathcal{I}_B$ is similar to Eq. (184b), but with A and B swapped: $\mathcal{P}_B = \mathcal{I}_B \Rightarrow \overline{\mathcal{P}}_A \succ \mathcal{P}_B = \overline{\mathcal{P}}_A \otimes \mathcal{I}_B$. As for the green zone, it is untouched as well, but the blue zone also sees its left quadrant shrunk into the top-left segment, so the green and blue zones become equivalent. This is Eq. (181a).

When $\overline{\mathcal{P}}_A = \mathcal{D}_A$ and $\mathcal{P}_B = \mathcal{I}_B$ (Fig. 10b), the yellow and green zones are shrunk to the segment $D \otimes B$. The versions of Eqs. (184a) and (184b) where A and B have swapped roles are simultaneously satisfied, so the diagram depicting the no signaling transformation (yellow) is equivalent to the one depicting the transformation which is one-way signaling to A (green). At the same time, the pink and blue zones are reduced to $A \otimes B - A \otimes D$. Eqs. (181a) and (181b) are simultaneously satisfied so the diagram depicting the two-way signaling transformation (blue) is equivalent to the one depicting the transformation which is one-way signaling to B (pink).

C No signaling is local quasi-orthogonality

Here we give more details about how the generalization of a quasi-orthogonality, Eq. (17), into party-wise quasi-orthogonality, Eq. (48), yields the condition (43). In order to justify its denomination of “no signaling”, we also show how

this condition relates to the model-independent notion of no signaling. That is, to a constraint on the joint outcome distribution that two parties observe when they act probabilistically on a state structure, i.e., when applying resolution of functionals locally.

Translating this statement into mathematical conditions, if Alice chooses her measurement settings depending on a classical random variable x and gets classical outcome a and if Bob does the same with settings y and outcome b we want the following constraints on their correlations:

$$\forall y, y', \quad \sum_b p(a, b|x, y) = \sum_b p(a, b|x, y') ; \quad (190a)$$

$$\forall x, x', \quad \sum_a p(a, b|x, y) = \sum_a p(a, b|x', y) . \quad (190b)$$

This is the standard statement that no local measurement scheme can be used to deterministically gain knowledge of the other party’s actions. The first condition, Eq. (190a), states that Alice’s distribution of outcome a cannot be used to determine which setting y Bob has used, thus that the measurement result of Alice cannot be used to guess Bob’s own choice of measurement. These correlations are *no signaling* from Bob to Alice. Respectively, the second condition, Eq. (190b), defines no signaling from Alice to Bob correlations.

For formulating this statement in the formalism developed in this article, assume two parties Alice and Bob trying to measure a shared bipartite state structure. Each locally sees its own measurement as part of a functional state structure, respectively $\overline{\mathcal{A}}$ and $\overline{\mathcal{B}}$. Alice’s and Bob’s measurements are locally disconnected and thus assumed in tensor product (otherwise, they can be trivially signaling to one another just by using the non-commutative nature of their operations). Thus, their joint operations belong to a state structure $\overline{\mathcal{X}}$ having the inner structure $\overline{\mathcal{X}} \subseteq \overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$. The question reduces to finding what kind of shared state W Alice and Bob can locally measure so that their outcome distributions are no signaling.

Let W be a state shared by Alice and Bob, $W \in \overline{\mathcal{X}} \subseteq \overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$. A measurement of the part of the state that is under the control of Alice is represented by a collection of operators $\{N_{a|x}\}$

resolving some operator $N_{|x} \in \overline{\mathcal{A}}$. Each $N_{a|x}$ is thus associated with Alice seeing outcome ‘ a ’ given setting ‘ x ’. The ‘ $|x$ ’ subscript notation is here to emphasize that x can condition the choice of the resolved operator as well as the choice of how it splits into effects, depending on Alice’s choice of strategy. The same way, Bob’s measurement is given by $\{N_{b|y}\}$. The probability distribution reads

$$p(a, b|x, y) = \text{Tr} \left[(N_{a|x} \otimes N_{b|y}) \cdot W \right], \quad (191)$$

which translates the no signaling conditions (190) into

$$\begin{aligned} \forall y, y', \\ \text{Tr} \left[(N_{a|x} \otimes N_{|y}) \cdot W \right] &= \text{Tr} \left[(N_{a|x} \otimes N_{|y'}) \cdot W \right]; \end{aligned} \quad (192a)$$

$$\begin{aligned} \forall x, x', \\ \text{Tr} \left[(N_{|x} \otimes N_{b|y}) \cdot W \right] &= \text{Tr} \left[(N_{|x'} \otimes N_{b|y}) \cdot W \right]. \end{aligned} \quad (192b)$$

These equations can be reduced into a more concise condition that resembles Eq. (17): take Eq. (192a), since it should hold for any y, y' the choice of a particular setting is no longer needed, and one can simply consider different $N \in \overline{\mathcal{B}}$. Rewriting it as

$$\begin{aligned} \forall N, N', \\ \text{Tr}_A \left[N_{a|x} \cdot \text{Tr}_B [(\mathbb{1} \otimes N) \cdot W] \right] &= \\ \text{Tr}_A \left[N_{a|x} \cdot \text{Tr}_B [(\mathbb{1} \otimes N') \cdot W] \right], \end{aligned} \quad (193)$$

one can simplify further by noticing that the possible $N_{a|x}$ actually range over the whole of $\mathcal{L}(\mathcal{H}^A)$ (up to trace normalization of course), so that only the trace over B part is relevant:

$$\forall N, N', \text{Tr}_B [(\mathbb{1} \otimes N) \cdot W] = \text{Tr}_B [(\mathbb{1} \otimes N') \cdot W]. \quad (194)$$

Finally, as $\mathbb{1}/c_B$ is a valid element of $\overline{\mathcal{B}}$, we can replace N' by it to obtain the shortened

$$\forall N, \text{Tr}_B [(\mathbb{1} \otimes N) \cdot W] = \frac{1}{c_B} \text{Tr}_B [(\mathbb{1} \otimes \mathbb{1}) \cdot W]. \quad (195)$$

Using Proposition 1, observe that $1/c_B = \text{Tr}[N]/d_B$ because of Eq. (15b), we can thus make $\text{Tr}[N]$ appear in the right-hand side and reach the desired formulation. Doing the same

reasoning for condition (190b), we obtain the following rephrasing of Eqs. (190):

$$\forall N_B, \text{Tr}_B [(\mathbb{1} \otimes N_B) \cdot W] = \frac{\text{Tr}[N_B] \text{Tr}_B[W]}{d_B}; \quad (196a)$$

$$\forall N_A, \text{Tr}_A [(N_A \otimes \mathbb{1}) \cdot W] = \frac{\text{Tr}[N_A] \text{Tr}_A[W]}{d_A}, \quad (196b)$$

Hence, no signaling (190) can be recast into the conditions (196), which amounts to requiring quasi-orthogonality for only one of the tensor factors of an operator. This leads to Eq. (47) and the ensuing characterization in Lemma 5, and in particular to Eqs. (61) and Corollary 2 in the main text.

D Proofs

D.1 Proof of Proposition 1

As $c_A/d_A \mathbb{1}$ is a valid element of $\mathcal{A} \subseteq \mathcal{L}(\mathcal{H}^A)$, any element N of $\overline{\mathcal{A}}$ must satisfy $1 = \langle N, c_A/d_A \mathbb{1} \rangle = \text{Tr}[N \cdot c_A/d_A \mathbb{1}] = c_A/d_A \text{Tr}[N]$. Since N is arbitrary, this fixes the normalization for all elements of $\overline{\mathcal{A}}$: $c_{\overline{\mathcal{A}}} \equiv d_A/c_A = \text{Tr}[N]$; this is condition (15b).

Positivity follows from the requirement of element-wise positivity with each resolution. For any N , let $\{|e_j\rangle\}$ is an orthonormal basis of $\mathcal{H}^A$ so that N is diagonal in this basis. Then its eigendecomposition is given by $N = \sum_j a_j |e_j\rangle\langle e_j|$, where $a_j \in \mathbb{C}$ are the eigenvalues of N . Since $c_A/d_A \mathbb{1} = c_A/d_A \sum_j |e_j\rangle\langle e_j|$, a resolution of $c_A/d_A \mathbb{1}$ can be defined as $\{E_j = c_A/d_A |e_j\rangle\langle e_j|\}_{j=0}^{d_A-1}$. That each element in the resolution gives a number between 0 and 1 reads $\forall j, \langle N, E_j \rangle \in [0, 1]$. But $\langle N, E_j \rangle = \text{Tr}[(\sum_i a_i |e_i\rangle\langle e_i|)^\dagger \cdot (c_A/d_A |e_j\rangle\langle e_j|)] = c_A/d_A \sum_i a_i^* \delta_{i,j}$, where $*$ indicates complex conjugation and $\delta_{i,j}$ is the Kronecker delta. Hence, $\forall j, \langle N, E_j \rangle = c_A/d_A a_j^* \in [0, 1]$, and since c_A/d_A is a positive real constant, this means that each a_j is real and greater or equal to zero, therefore that N is positive semi-definite, condition (15a).

Finally, the projective condition remains. First, $1/c_A \mathbb{1}$ must belong to $\overline{\mathcal{A}}$ as it is positive, properly normalized, and respects $\text{Tr}[1/c_A \mathbb{1} \cdot V] = 1$ for all $V \in \mathcal{A}$. Assuming that $\overline{\mathcal{A}}$ is a deterministic state structure

with projector $\bar{\mathcal{P}}_A \equiv \mathcal{I} - \mathcal{P}_A + \mathcal{D}$, the if part follows: Any positive and properly normalized operator N in $\mathcal{L}(\mathcal{H}^A)$ on which the projector $\bar{\mathcal{P}}_A$ is applied obeys $\langle \bar{\mathcal{P}}_A \{N\}, V \rangle = 1$ for all $V \in \mathcal{A}$ because $\langle \bar{\mathcal{P}}_A \{N\}, V \rangle \equiv \langle (\mathcal{I} - \mathcal{P}_A) \{N\}, V \rangle + \langle \mathcal{D}\{N\}, V \rangle$. The first member on the right part of the equality vanishes because it belongs to the orthogonal complement of $\mathcal{A}$; the second member is normalized so that $\langle \mathcal{D}\{N\}, V \rangle = \langle \text{Tr}[N]/d_A \mathbb{1}, V \rangle = \langle \frac{d_A}{c_A d_A} \mathbb{1}, V \rangle = 1/c_A \text{Tr}[V] = 1$.

To prove the only if part, assume that there is a positive X with trace norm $c_{\bar{A}} = d_A/c_A$ that gives $\langle X, V \rangle = 1$ for all $V \in \mathcal{A}$. Then, $\langle X, \mathcal{P}_A \{V\} \rangle = 1$ by definition of V . Since $\mathcal{P}_A$ is self-adjoint, this further implies $\langle \mathcal{P}_A \{X\}, V \rangle = 1$. Remark that $\mathcal{P}_A = \mathcal{D} + \mathcal{P}_A - \mathcal{D}$ and that $\mathcal{P}_A - \mathcal{D} = \mathcal{I} - \bar{\mathcal{P}}_A$ so $\mathcal{P}_A = \mathcal{D} + (\mathcal{I} - \bar{\mathcal{P}}_A)$. Inserting this in the inner product and using its linearity, it must be that $\langle \mathcal{D}\{X\}, V \rangle + \langle (\mathcal{I} - \bar{\mathcal{P}}_A)\{X\}, V \rangle = 1$. The first term of the l.h.s. simplifies into $\langle \mathcal{D}\{X\}, V \rangle = d_A/c_A \langle \frac{1}{d_A}, V \rangle$ because the trace of X was assumed to be d_A/c_A and so this term must be equal to 1 as V has a trace norm of c_A by definition. Therefore, it must be the case that $\langle (\mathcal{I} - \bar{\mathcal{P}}_A)\{X\}, V \rangle = 0$. Since both X and V are generally non-zero (as they are arbitrary), and we want this inner product to be zero for all V , there are two possibilities for X : it can either belong to the orthogonal complement of V , i.e. to the image of $\mathcal{I} - \mathcal{P}_A$, or else it can be that $(\mathcal{I} - \bar{\mathcal{P}}_A)\{X\} = 0$, i.e. that $\bar{\mathcal{P}}_A\{X\} = X$. Because X has a non-zero trace while the orthogonal complement of $\mathcal{P}_A$ is spanned by traceless matrices, we conclude that only the latter condition can be true, i.e. $\bar{\mathcal{P}}_A\{X\} = X$ if $\langle X, V \rangle = 1$ for all $V \in \mathcal{A}$. This proves the necessity of the projective condition.

D.2 Proof of Proposition 2

Following Ref. [24], the second requirement of Definition 6 states that:

$$\mathcal{M}(V) \geq 0, \quad (197a)$$

$$\text{Tr}[\mathcal{M}(V)] = c_B, \quad (197b)$$

$$\mathcal{P}_B \circ \mathcal{M} \circ \mathcal{P}_A = \mathcal{M} \circ \mathcal{P}_A. \quad (197c)$$

The third requirement imposes condition (23a). This condition implies in turn the positivity of

$\mathcal{M}$, Eq. (197a), because $\mathcal{M}$ having a positive Choi matrix is equivalent to $\mathcal{M}$ being completely positive by Choi theorem [16].

Using the reverse definition of the CJ isomorphism, Eq. (2), condition (197b) becomes (subscripts have been put on M and V for clarity)

$$\text{Tr}[M_{AB} \cdot (V_A \otimes \mathbb{1}_B)] = c_B, \quad (198)$$

from which the normalization condition (23b) follows, as V can be $c_A \mathbb{1}/d_A$. Finally, as $\mathcal{M}(V) = (\text{Tr}_A[M \cdot (V \otimes \mathbb{1}_B)])^T$ and $\mathcal{M}(V) \in \mathcal{B}$, it should be true that, $\forall N \in \mathcal{B}$,

$$\begin{aligned} 1 &= \text{Tr}_B[N \cdot \mathcal{M}(V)] \\ &= \text{Tr}_B[N \cdot (\text{Tr}_A[M \cdot (V \otimes \mathbb{1}_B)])^T] \\ &= \left(\text{Tr}_B[\text{Tr}_A[M \cdot (V \otimes \mathbb{1}_B)] \cdot N^T] \right)^T \quad (199) \\ &= \text{Tr}[M \cdot (V \otimes N^T)]. \end{aligned}$$

Since M is positive and normalized, and since $V \in \mathcal{A}$ and $N \in \mathcal{B}$ are arbitrary, it follows from this last equation that the set of all M obeying the above is a functional state structure on a spanning set for $\mathcal{A} \otimes \bar{\mathcal{B}}^T \cong \mathcal{A} \otimes \bar{\mathcal{B}}$ (recall that a state structure is closed under transposition). From linearity, it means that M is also normalized on affine combinations of this set. Indeed, let $\{V_i\}$ and $\{N_i\}$ two arbitrary sets of states in $\mathcal{A}$ and $\bar{\mathcal{B}}$, respectively, and let $1 = \sum_i q_i$ where $\forall i, q_i \in \mathbb{R}$. Then,

$$\begin{aligned} 1 &= \sum_i q_i \\ &= \sum_i q_i \text{Tr}[M \cdot (V_i \otimes N_i^T)] \quad (200) \\ &= \text{Tr}\left[M \cdot \left(\sum_i q_i (V_i \otimes N_i^T)\right)\right]. \end{aligned}$$

Hence, M is also normalized on the affine hull of the pure tensor products of $\mathcal{A}$ and $\bar{\mathcal{B}}$, which is $\mathcal{A} \otimes \bar{\mathcal{B}}$ by Definition 5. Therefore, using Proposition 1, the set of positive and normalized operators on this set is characterized by a projector $\bar{\mathcal{P}}_A \otimes \bar{\mathcal{P}}_B$ yielding the projective condition (23c) (note this was first proven in [5, Proposition 1]).

D.3 Proof of Proposition 3

The characterization of the set of deterministic events is one of the core results of Ref. [5]. The part of their derivation that is of interest for this

proof is the characterization of what is called the *deterministic events* of an arbitrary type X constructed from elementary types $A, B, \dots, K$ and using the $\{1, (,), \rightarrow\}$ connectors. The authors start by 1) defining the deterministic events of elementary type (Definition 6, p.6); then they 2) generalize the definition to deterministic events of arbitrary type $X \rightarrow Y$, where X and Y can be any type (Definition 7, p. 7); and finally they 3) characterize the deterministic events of type $X \rightarrow Y$ (Proposition 1, p. 11). The end goal of this proof is then to show that their Proposition 1 is our Proposition 2 in the special case where the base state structures are sets of density matrices.

Our proof strategy will consist of showing that each of these steps has an equivalent in the projective characterization under the assumption that the base state structures are the quantum states. Precisely, we will show 1) what is called ‘deterministic events of elementary types $A, \dots, K$ ’, and noted $T_1(A), \dots, T_1(K)$ in Ref. [5], correspond to the base state structures $\mathcal{A}, \dots, \mathcal{K}$, all associated with the identity projector, $\mathcal{P}_A = \mathcal{I}_A, \dots, \mathcal{P}_K = \mathcal{I}_K$, and with unit normalization, $c_A = 1, \dots, c_K = 1$; 2) that what is called ‘the sets of deterministic events of type X obtained from elementary types $A, \dots, K$ ’, and noted $T_1(X)$ in Ref. [5] (remark that for non-elementary types, the authors of this work prefer to use small case letters like x , but this convention is not followed in this article) corresponds to a state structure $\mathcal{X}$ obtained from base structures $\mathcal{A}, \dots, \mathcal{K}$; 3) that their characterization of the deterministic events of type $X \rightarrow Y$ corresponds to our characterization of the state structure $\mathcal{X} \rightarrow \mathcal{Y}$.

The first item is trivial to show. By Definition 6 in Ref. [5], an operator M is a deterministic event of elementary type A if $M \in \mathcal{L}(\mathcal{H}^A) : M \geq 0$ and $\text{Tr}[M] = 1$. This is exactly the state structure of density operators as characterized by Eqs. (8).

The second item is shown by first using Theorem 3 of Ref. [5], which characterizes the set of deterministic events of arbitrary type by the following constraints:

$$M \in T_1(X \rightarrow Y) \iff M \geq 0; \quad (201a)$$

$$\forall V \in T_1(X) : \mathcal{M}(V) \in T_1(Y); \quad (201b)$$

where $\mathcal{M}$ is the linear map obtained when apply-

ing the reverse direction of the CJ correspondence on M . In such a form, it is direct to see that it corresponds to our definition of a structure-preserving map, Def. 6.

For the third item, Proposition 1 of Ref. [5] provides the following characterization result:

$$M \in T_1(X \rightarrow Y) \iff M \geq 0, \quad (202a)$$

$$M = \lambda_{X \rightarrow Y} \mathbb{1} + H_{X \rightarrow Y}, \quad (202b)$$

where $H_{X \rightarrow Y}$ belongs to a subspace $\Delta_{X \rightarrow Y}$. This subspace and the constant $\lambda_{X \rightarrow Y}$ are defined recursively by:

For $X = 1$ and Y elementary,

$$\lambda_Y = \frac{1}{d_Y}, \quad (203a)$$

$$\Delta_Y = \{M \in \mathcal{L}(\mathcal{H}^Y) \mid \text{Tr}[M] = 0\}. \quad (203b)$$

Else,

$$\lambda_{X \rightarrow Y} = \frac{\lambda_Y}{d_X \lambda_X}, \quad (203c)$$

$$\Delta_{X \rightarrow Y} = [\mathcal{L}(\mathcal{H}^X) \otimes \Delta_Y] \oplus [\overline{\Delta}_X \otimes \Delta_Y^\perp]. \quad (203d)$$

In the last formula, the $\overline{}$ indicates a quasi-orthogonal complement whilst the $^\perp$ indicates an orthogonal complement.

What we want to show is that the conditions (202) are exactly the conditions (23), i.e.,

$$M \in \mathcal{X} \rightarrow \mathcal{Y} \iff M \geq 0, \quad (204a)$$

$$\text{Tr}[M] = \frac{c_Y}{c_X} d_X, \quad (204b)$$

$$(\mathcal{P}_X \rightarrow \mathcal{P}_Y) \{M\} = M, \quad (204c)$$

despite our vastly different approaches. Firstly, notice that conditions (202a) and (204a) are the same, so the proof will consist in showing that condition (202b) is equivalent to conditions (204b) and (204c).

Since any Δ_X space is traceless by construction (see Sec. 5. b. of Ref. [5]), we can extract the trace condition by taking the trace of Eq.(202b), yielding $\text{Tr}[M] = 1$ for Eq. (203a), and $\text{Tr}[M] = \frac{\lambda_Y}{\lambda_X} d_Y$ for Eq. (203c). By direct inspection, one can identify that $c_X = \lambda_X d_X$, so that $\text{Tr}[M] = \frac{c_Y}{c_X} d_X$, recovering Eq. (204b), and that $c_Y = 1$ when $\mathcal{Y}$ is a base state structure, i.e. when $\mathcal{P}_Y = \mathcal{I}_Y$.

We can characterize the Δ subspaces using superoperator projectors: Eq. (203b) is equivalent to the condition $(\mathcal{I} - \mathcal{D})\{M\} = M$ since the trace of an operator is contained in the subspace invariant under the depolarizing superoperator $\mathcal{D}$. Injecting this in Eq. (202b), we recover the (trivial) condition $\mathcal{I}\{M\} = M$ characterizing the set of density matrices in the case where $X = 1$ and Y is elementary.

Now, for the general case, since Δ_X (with X being any type) must be a traceless subspace, we can posit it to correspond to some projector $\mathcal{P}_X - \mathcal{D}_X$ such that $V \in \Delta_X \iff (\mathcal{P}_X - \mathcal{D}_X)\{V\} = V$. Similarly, to Δ_Y will correspond some projector $\mathcal{P}_Y - \mathcal{D}_Y$ and to $\Delta_{X \rightarrow Y}$ some $\mathcal{P}_{X \rightarrow Y} - \mathcal{D}_X \otimes \mathcal{D}_Y$. Injecting the three ansätze into Eq. (203d) yields

$$\begin{aligned} \mathcal{P}_{X \rightarrow Y} - \mathcal{D}_X \otimes \mathcal{D}_Y &= [\mathcal{I}_X \otimes (\mathcal{P}_Y - \mathcal{D}_Y)] \\ &+ [(\overline{\mathcal{P}_X - \mathcal{D}_X}) \otimes (\mathcal{P}_Y - \mathcal{D}_Y)^\perp]. \end{aligned} \quad (205)$$

(Note that the correspondence between the direct sum of spaces $\oplus$ and the sum of projectors is only possible because the intersection of each part of the sum is the zero projector. In general, the union of spaces through $\oplus$ must correspond to the union of projectors $\cup$ introduced in Sec. 4.)

The orthogonal complement $\mathcal{P}^\perp$ of a projector $\mathcal{P}$ is given by $\mathcal{I} - \mathcal{P}$ whereas his quasi-orthogonal complement $\overline{\mathcal{P}}$ is given by $\mathcal{I} - \mathcal{P} + \mathcal{D}$ (see Eq. (15c)). Injecting these into the right-hand side, then subsequently simplifying, gives

$$\begin{aligned} &[\mathcal{I}_X \otimes (\mathcal{P}_Y - \mathcal{D}_Y)] \\ &+ [(\overline{\mathcal{P}_X - \mathcal{D}_X}) \otimes (\mathcal{P}_Y - \mathcal{D}_Y)^\perp] \\ &= [\mathcal{I}_X \otimes (\mathcal{P}_Y - \mathcal{D}_Y)] + [(\mathcal{I}_X - \mathcal{P}_X) \otimes \overline{\mathcal{P}_Y}] \\ &= \mathcal{I}_X \otimes \mathcal{P}_Y - \mathcal{I}_X \otimes \mathcal{D}_Y + \mathcal{I}_X \otimes \overline{\mathcal{P}_Y} - \mathcal{D}_X \otimes \overline{\mathcal{P}_Y} \\ &= \mathcal{I}_X \otimes \mathcal{I}_Y - \mathcal{P}_X \otimes \overline{\mathcal{P}_Y}. \end{aligned} \quad (206)$$

(All the identities used and projector manipulations are detailed in Sec. 4 and App. 84.) Hence, Eq. (203d) is recast as the following projector identity:

$$\mathcal{P}_{X \rightarrow Y} - \mathcal{D}_X \otimes \mathcal{D}_Y = \mathcal{I}_X \otimes \mathcal{I}_Y - \mathcal{P}_X \otimes \overline{\mathcal{P}_Y}. \quad (207)$$

Passing the $\mathcal{D}_X \otimes \mathcal{D}_Y$ term on the other side, we obtain

$$\begin{aligned} \mathcal{P}_{X \rightarrow Y} &= \mathcal{I}_X \otimes \mathcal{I}_Y - \mathcal{P}_X \otimes \overline{\mathcal{P}_Y} + \mathcal{D}_X \otimes \mathcal{D}_Y \\ &= \overline{\mathcal{P}_X \otimes \overline{\mathcal{P}_Y}} \\ &= \mathcal{P}_X \rightarrow \mathcal{P}_Y, \end{aligned} \quad (208)$$

recovering Eq. (204c).

Therefore, we showed that conditions (202) are equivalent to the conditions (204). We conclude that Proposition 1 of Ref. [5] is equivalent to Proposition 2 of this paper. This concludes the proof.

D.4 Proof of Lemma 3

Let $\mathcal{P}^{(n)}$ be the projector obtained after n ‘steps’ that can be categorized as 1) do a global negation of the projector, $\mathcal{P}^{(n)} = \overline{\mathcal{P}^{(n-1)}}$; 2) add a base projector on the right using tensor product, $\mathcal{P}^{(n)} = \mathcal{P}^{(n-1)} \otimes \mathcal{P}_X$; 3) add a similarly obtained projector after k steps, $\mathcal{P}^{(k)}$, on the right using tensor product, $\mathcal{P}^{(n)} = \mathcal{P}^{(n-1)} \otimes \mathcal{P}^{(k)}$. This covers all cases as the $\rightarrow$ can be split into a negation and a tensor, $\cdot \rightarrow \cdot \equiv \cdot \otimes \overline{\cdot}$, and one can always redefine the tensor factors labeling so that the added system is on the right since $\mathcal{H}^A \otimes \mathcal{H}^B \cong \mathcal{H}^B \otimes \mathcal{H}^A$.

The only non-trivial first step is to choose a base projector, say $\mathcal{P}_A$, in which case the claim trivially holds, $\mathcal{P}_A \subseteq \mathcal{P}_A \subseteq \overline{\mathcal{P}_A}$. Suppose it holds after $n - 1$ steps during which a certain amount of base projectors were added so that the joint projectors act on a Hilbert space whose factors range from $\mathcal{H}^A$ up to $\mathcal{H}^J$ (this is without loss of generality as the index of the last factor, J , was arbitrarily chosen). Then, the hypothesis for the induction reads

$$\tilde{\mathcal{P}}_A \otimes \dots \otimes \tilde{\mathcal{P}}_J \subseteq \mathcal{P}^{(n-1)} \subseteq \overline{\overline{\mathcal{P}_A} \otimes \dots \otimes \overline{\mathcal{P}_J}}. \quad (209)$$

Let $\tilde{\mathcal{P}}_A \otimes \dots \otimes \tilde{\mathcal{P}}_J \equiv \mathcal{P}_{NS}^{(n-1)}$ and $\overline{\overline{\mathcal{P}_A} \otimes \dots \otimes \overline{\mathcal{P}_J}} \equiv \mathcal{P}_{FS}^{(n-1)}$. Then, $\mathcal{P}_{NS}^{(n-1)} \subseteq \mathcal{P}^{(n-1)} \subseteq \mathcal{P}_{FS}^{(n-1)}$. The following holds because of Eqs. (31)

$$\overline{\mathcal{P}_{NS}^{(n-1)}} \supseteq \overline{\mathcal{P}^{(n-1)}} \supseteq \overline{\mathcal{P}_{FS}^{(n-1)}}, \quad (210)$$

This corresponds to doing a step of category 1), $\mathcal{P}^{(n)} = \overline{\mathcal{P}^{(n-1)}}$, in which case $\overline{\mathcal{P}_{NS}^{(n-1)}} = \overline{\mathcal{P}_A} \otimes \dots \otimes \overline{\mathcal{P}_J}$. A negation can be put over every single projector by redefining the tilde, $\tilde{\mathcal{P}}_A \mapsto \overline{\mathcal{P}_A}$, which implies that $\overline{\mathcal{P}_{NS}^{(n-1)}}$ is redefined as $\overline{\mathcal{P}_{NS}^{(n-1)}} = \mathcal{P}_{FS}^{(n)} = \overline{\overline{\mathcal{P}_A} \otimes \dots \otimes \overline{\mathcal{P}_J}}$. The same can be done on $\overline{\mathcal{P}_{FS}^{(n-1)}}$, which proves the induction for 1).

Let $\mathcal{P}'$ be an arbitrary projector on operator system, the following holds by (136) and (33)

proven in App. A.4.1:

$$\mathcal{P}_{NS}^{(n-1)} \otimes \mathcal{P}' \subseteq \mathcal{P}^{(n-1)} \otimes \mathcal{P}' \subseteq \mathcal{P}_{FS}^{(n-1)} \otimes \mathcal{P}', \quad (211a)$$

$$\mathcal{P}_{FS}^{(n-1)} \otimes \mathcal{P}' \subseteq \overline{\mathcal{P}_{FS}^{(n-1)}} \otimes \overline{\mathcal{P}'}, \quad (211b)$$

$$\mathcal{P}_{NS}^{(n-1)} \otimes \mathcal{P}' \subseteq \mathcal{P}^{(n-1)} \otimes \mathcal{P}' \subseteq \overline{\mathcal{P}_{FS}^{(n-1)}} \otimes \overline{\mathcal{P}'}. \quad (211c)$$

Since $\mathcal{P}'$ is arbitrary, the first equation corresponds to either doing step 2) or 3). For case 2), $\mathcal{P}' \equiv \mathcal{P}_L$ is the $(j+1)$ -th subsystem added, corresponding to some party L , so that $\mathcal{P}^{(n-1)} \otimes \mathcal{P}_L = \mathcal{P}^{(n)}$. The third equation then reads $\mathcal{P}_{NS}^{(n-1)} \otimes \mathcal{P}_L \subseteq \mathcal{P}^{(n)} \subseteq \overline{\mathcal{P}_{FS}^{(n-1)}} \otimes \overline{\mathcal{P}'}$. Set $\mathcal{P}_L \equiv \tilde{\mathcal{P}}_L$ and the induction is proven. The reasoning is analog in case 3). Note that the added system, $\mathcal{P}' \equiv \mathcal{P}^{(k)}$, is also included in some $\mathcal{P}_{NS}^{(k)} \subseteq \mathcal{P}^{(k)} \subseteq \mathcal{P}_{FS}^{(k)}$ by assumption. Using Eq. (136) again, one has $\mathcal{P}_{NS}^{(n-1)} \otimes \mathcal{P}_{NS}^{(k)} \subseteq \mathcal{P}^{(n-1)} \otimes \mathcal{P}^{(k)} \subseteq \overline{\mathcal{P}_{FS}^{(n-1)}} \otimes \overline{\mathcal{P}_{FS}^{(k)}}$. As $\mathcal{P}^{(n)} = \mathcal{P}^{(n-1)} \otimes \mathcal{P}^{(k)}$, the induction is proven by using the fact that the tensor as well as $\overline{\cdot} \rightarrow \cdot$ are associative operations to extend the expressions on both side of $\mathcal{P}^{(n-1)} \otimes \mathcal{P}^{(k)}$ and then to define the $\tilde{\mathcal{P}}$'s according to the sought expression.

D.5 Proof of Lemma 5

Let $\{\sigma_i^X\}$ be a basis of $\mathcal{L}(\mathcal{H}^X)$ so that

$$(\sigma_i^X)^\dagger = \sigma_i^X; \quad (212a)$$

$$\sigma_0^X \equiv \mathbb{1}/\sqrt{d_A}; \quad (212b)$$

$$\text{Tr}[\sigma_{i \neq 0}^X] = 0; \quad (212c)$$

$$\langle \sigma_i^X, \sigma_j^X \rangle \equiv \text{Tr}[\sigma_i^X \cdot \sigma_j^X] = \delta_{i,j}. \quad (212d)$$

And choose this basis on $\mathcal{L}(\mathcal{H}^A)$ such that the first $n < d_A^2$ elements after σ_0^A , i.e. the set $\{\sigma_1^A, \sigma_2^A, \dots, \sigma_n^A\}$, form a spanning set of $\mathcal{A} \setminus \{\mathbb{1} = \sqrt{d_A} \sigma_0\}$. This implies that if $\mathcal{P}_A$ is the projector on $\mathcal{A}$, then $\sum_{i=0}^{d_A^2-1} \mathcal{P}_A \{\sigma_i^A\} = \sum_{i=0}^n \mathcal{P}_A \{\sigma_i^A\}$. We expand M and V using this basis, let $M = \sum_{i,j}^{d_A^2-1, d_B^2-1} m_{ij} \sigma_i^A \otimes \sigma_j^B$; $V = c_A/\sqrt{d_A} \sigma_0^A + \sum_{k=1}^n v_k \sigma_k^A$. By direct computation:

$$\text{Tr}_A[M] = d_A \sum_{j=0}^{d_B^2-1} m_{0j} \sigma_j^B; \quad (213)$$

$$\text{Tr}[V] = c_A; \quad (214)$$

$$\text{Tr}_A[M \cdot (V \otimes N)] = \left(\sum_{i=0}^n \sum_{j=0}^{d_B^2-1} d_A m_{ij} v_i \sigma_j^B \right) \cdot N. \quad (215)$$

Combining Eqs. (213) and (214) in the right-hand side of Eq. (48) and (215) in its left-hand side yields

$$\begin{aligned} \text{Tr}_A[M \cdot (V \otimes N)] &= \frac{\text{Tr}_A[M] \text{Tr}[V]}{d_A} \cdot N \Rightarrow \\ &= \left(\sum_{i=0}^n \sum_{j=0}^{d_B^2-1} d_A m_{ij} v_i \sigma_j^B \right) \cdot N = \\ &= \frac{\left(d_A \sum_{j=0}^{d_B^2-1} m_{0j} \sigma_j^B \right) (c_A)}{d_A} \cdot N, \quad (216) \end{aligned}$$

the equality is true for

$$\left(\sum_{i>0}^n \sum_{j=0}^{d_B^2-1} m_{ij} v_i \sigma_j^B \right) \cdot N = 0. \quad (217)$$

This reduces to $\sum_{i>0}^n \sum_{j=0}^{d_B^2-1} m_{ij} v_i \sigma_j^B = 0$ because N is arbitrary and can therefore be taken as $\mathbb{1}_B$. On the other hand, this equation has to hold for all V in $\mathcal{A}$, so one can consider the family $V_i \equiv \sigma_i^A$. Each V_i then induces a condition $\sum_j m_{ij} \sigma_j^B = 0$ for $i \in 1, \dots, n$. Because the σ_j^B 's are orthonormal vectors, and M is arbitrary, this condition is true for a given V_i if and only if $m_{ij} = 0 \forall j$ whenever i corresponds to a basis element of $\mathcal{A}$, i.e., whenever $0 < i \leq n$. The condition for the equality to hold is therefore that M has the form $\sum_{j=0}^{d_B^2-1} \left(m_{0j} \sigma_0^A \otimes \sigma_j^B + \sum_{i>n}^{d_A^2-1} m_{ij} \sigma_i^A \otimes \sigma_j^B \right)$. This is equivalent to requiring Eq. (49) on M , concluding the proof.

D.6 Proof of Proposition 6

General case: That Γ is a projector on operator system acting on $\mathcal{L}(\mathcal{H}^A \otimes \dots \otimes \mathcal{H}^K)$ is direct from the properties of the operations in the algebra proven in App. A. Note that these operations also guarantee commutativity, so that the composition (cap) of any two such Γ is always well-defined and results in a projector on an operator system. As a consequence, the cup is also well-defined. The least and greatest elements directly follow as any expression must be contained between the projector on the span

of the identity $\mathcal{D}_{A\dots K}$ and the identity projector $\mathcal{I}_{A\dots K}$, so that $\mathcal{D}_{A\dots K} \stackrel{(128a)}{=} \mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K \subseteq \Gamma \subseteq \mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K \stackrel{(128b)}{=} \mathcal{I}_{A\dots K}$. Finally, the least and greatest elements are reached since (for example) one can always consider the elementary projector $\mathcal{P}_A \otimes \dots \otimes \mathcal{P}_K$ and its party-wise-negated counterpart $\overline{\mathcal{P}}_A \otimes \dots \otimes \overline{\mathcal{P}}_K$. These two projectors are valid Γ 's by construction, and their intersection and union yield $\mathcal{D}_A \otimes \dots \otimes \mathcal{D}_K$ and $\mathcal{I}_A \otimes \dots \otimes \mathcal{I}_K$, respectively.

Restricted case: This is essentially proven the same way as Lemma 3, see App. D.4, but with two new categories of ‘steps’: 4) $\mathcal{P}^{(n)} = \mathcal{P}^{(n-1)} \prec \mathcal{P}_L$ and 5) $\mathcal{P}^{(n)} = \mathcal{P}^{(n-1)} \prec \mathcal{P}^{(k)}$. Reducing these steps into a chain of inclusion like $\mathcal{P}_{NS}^{(n-1)} \otimes \mathcal{P}' \subseteq \mathcal{P}^{(n-1)} \prec \mathcal{P}' \subseteq \overline{\mathcal{P}_{NS}^{(n-1)}} \otimes \overline{\mathcal{P}'}$ is again obtained by first noticing that $\mathcal{P}_A \subseteq \mathcal{P}'_A \subseteq \mathcal{P}''_A \Rightarrow \mathcal{P}_A \prec \mathcal{P}_B \subseteq \mathcal{P}'_A \prec \mathcal{P}_B \subseteq \mathcal{P}''_A \prec \mathcal{P}_B$ and then by using relations (60) so that $\mathcal{P}_A \otimes \mathcal{P}_B \subseteq \mathcal{P}_A \prec \mathcal{P}_B$ and $\mathcal{P}''_A \prec \mathcal{P}_B \subseteq \overline{\mathcal{P}''_A} \otimes \overline{\mathcal{P}_B}$.

D.7 Proof of Theorem 2

Let $\Gamma = \bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)} \right)$ and $\Xi = \bigcup_{m=1}^z \left(\bigcap_{n=1}^{t_m} \tilde{\mathcal{P}}_{\sigma_{mn}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{mn}(K)} \right)$ be two normal forms involving k base projectors. We define the shorthand notation: $\Gamma_i \equiv \left(\bigcap_{j=1}^{y_i} \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)} \right)$ so that

$$\Gamma = \bigcup_{i=1}^x \Gamma_i = \Gamma_1 \cup \Gamma_2 \cup \dots \cup \Gamma_x, \quad (218)$$

and $\Gamma_{ij} \equiv \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)}$ so that $\Gamma_i = \left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right)$,

$$\Gamma = \bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right). \quad (219)$$

Hence,

$$\begin{aligned} \Gamma &= \bigcup_{i=1}^x (\Gamma_{i1} \cap \Gamma_{i2} \cap \dots \cap \Gamma_{iy_i}) \\ &= \left(\bigcap_{j=1}^{y_1} \Gamma_{1j} \right) \cup \left(\bigcap_{j=1}^{y_2} \Gamma_{2j} \right) \cup \dots \cup \left(\bigcap_{j=1}^{y_x} \Gamma_{xj} \right); \end{aligned} \quad (220)$$

and Ξ_m and Ξ_{mn} are defined accordingly.

First, remark that the negation, intersection, and union of normal forms can be put into normal

forms: $\Gamma \cup \Xi$ is a normal form obtained simply by merging the unions:

$$\begin{aligned} \Gamma \cup \Xi &= \left(\bigcup_{i=1}^x \Gamma_i \right) \cup \left(\bigcup_{m=1}^z \Xi_m \right) \\ &= \Gamma_1 \cup \Gamma_2 \cup \dots \cup \Gamma_x \cup \Xi_1 \cup \Xi_2 \cup \dots \cup \Xi_z, \end{aligned} \quad (221)$$

which is a normal form once the redundant terms in the series –the Γ_i 's that happen to be equivalent to some Ξ_m 's– have been removed by commutativity and associativity. To prove that $\Gamma \cap \Xi$ can be put in normal form requires to use of the distribution law (105)

$$\begin{aligned} \Gamma \cap \Xi &= \bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right) \cap \bigcup_{m=1}^z \left(\bigcap_{n=1}^{t_m} \Xi_{mn} \right) \\ &= (\Gamma_1 \cup \Gamma_2 \cup \dots \cup \Gamma_x) \cap (\Xi_1 \cup \Xi_2 \cup \dots \cup \Xi_z) \\ &= (\Gamma_1 \cap \Xi_1) \cup (\Gamma_2 \cap \Xi_1) \cup \dots \cup (\Gamma_z \cap \Xi_1) \\ &\quad \cup (\Gamma_1 \cap \Xi_2) \cup \dots \cup (\Gamma_z \cap \Xi_t). \end{aligned} \quad (222)$$

Each $(\Gamma_i \cap \Xi_m)$ term is an intersection of intersections, therefore they can be merged as an overall intersection by associativity and the redundancy can be removed the same way as for the union case above. Then again, the unions of these terms is a normal form once the redundancies like $(\Gamma_i \cap \Xi_m) = (\Gamma_{i'} \cap \Xi_{m'})$ for some $(i, m) \neq (i', m')$ have been removed. To prove that $\bar{\Gamma}$ can be put in normal form requires to use successively the De Morgan laws,

$$\begin{aligned} \bar{\Gamma} &= \overline{\bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right)} \\ &\stackrel{(119a)}{=} \bigcap_{i=1}^x \overline{\left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right)} \\ &\stackrel{(119b)}{=} \bigcap_{i=1}^x \left(\bigcup_{j=1}^{y_i} \bar{\Gamma}_{ij} \right), \end{aligned} \quad (223)$$

then the commutation of the negation with the prec, Eq. (66) (proven at Eq. (162)), is used on each term,

$$\begin{aligned} \bar{\Gamma}_{ij} &= \overline{\tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)}} \\ &= \overline{\tilde{\mathcal{P}}_{\sigma_{ij}(A)}} \prec \dots \prec \overline{\tilde{\mathcal{P}}_{\sigma_{ij}(K)}} \\ &= \tilde{\mathcal{P}}'_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}'_{\sigma_{ij}(K)} = \Gamma'_{ij} \end{aligned} \quad (224)$$

where the base projectors have been redefined so as to incorporate the negation. Thus, the Γ'_{ij} are

normal forms, and so their intersections of unions can be put into normal form by using the two properties proven just before this one.

Next, let Υ be a projector on an operator system over l subsystems in $\mathcal{L}(\mathcal{H}^L \otimes \dots \otimes \mathcal{H}^R)$ that is in normal form, $\Upsilon = \bigcup_{m=1}^z \left(\bigcap_{n=1}^{t_m} \tilde{\mathcal{P}}_{\chi_{mn}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\chi_{mn}(R)} \right)$ where χ_{mn} is an element of the permutation group over l symbols. Then the one-way signaling composition $\Gamma \prec \Upsilon$ is a projector on operator system over $\mathcal{L}(\mathcal{H}^A \otimes \dots \otimes \mathcal{H}^K \otimes \mathcal{H}^L \otimes \dots \otimes \mathcal{H}^R)$ that can be put into a normal form as well. This is proven using the interchange laws (69). Define Υ_m and Υ_{mn} like above, let u ranging from 1 to $x \times z$ such that $u = 1$ is identified with $(i, m) = (1, 1)$, $u = 2$ with $(i, m) = (1, 2)$, etc., let v_u ranging from 1 to $y_i \times t_m$ such that $v_u = 1$ is identified with $(j, n) = (1, 1)$, etc., and let $\zeta_{uv} = (\sigma_{ij}, \chi_{mn})$ be an element of the permutation group over $k + l$ elements indexed by uv . That way, an element like $\Gamma_{ij} \prec \Upsilon_{mn}$ can be rewritten as Θ_{uv} in the following manner:

$$\begin{aligned} \Gamma_{ij} \prec \Upsilon_{mn} &= \left(\tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)} \right) \\ &\prec \left(\tilde{\mathcal{P}}_{\chi_{mn}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\chi_{mn}(R)} \right) \\ &= \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)} \\ &\prec \tilde{\mathcal{P}}_{\chi_{mn}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\chi_{mn}(R)} \\ &= \tilde{\mathcal{P}}_{\zeta_{uv}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\zeta_{uv}(K)} \\ &\prec \tilde{\mathcal{P}}_{\zeta_{uv}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\zeta_{uv}(R)} \\ &\equiv \Theta_{uv}; \quad (225) \end{aligned}$$

This is but a relabelling using associativity of the prec (67) (proven at Eq. (163)). But the same way, an element like $\Gamma_{ij} \succ \Upsilon_{mn}$ can also be identified with some $\Theta_{u'v'}$ by finding the permutation $\zeta_{u'v'}$ that exactly corresponds to $\tilde{\mathcal{P}}_{\zeta_{u'v'}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\zeta_{u'v'}(K)} \prec \tilde{\mathcal{P}}_{\zeta_{u'v'}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\zeta_{u'v'}(R)} = \tilde{\mathcal{P}}_{\chi_{mn}(L)} \prec \dots \prec \tilde{\mathcal{P}}_{\chi_{mn}(R)} \prec \tilde{\mathcal{P}}_{\sigma_{ij}(A)} \prec \dots \prec \tilde{\mathcal{P}}_{\sigma_{ij}(K)}$. The important thing to notice is by definition, Θ_{uv} is a normal form. The one-way sig-

naling composition of Γ with Υ then reads:

$$\begin{aligned} \Gamma \prec \Upsilon &= \bigcup_{i=1}^x \left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right) \prec \bigcup_{m=1}^z \left(\bigcap_{n=1}^{t_m} \Upsilon_{mn} \right) \\ &\stackrel{(173)}{=} \bigcup_{i=1}^x \left(\left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right) \prec \left(\bigcup_{m=1}^z \left(\bigcap_{n=1}^{t_m} \Upsilon_{mn} \right) \right) \right) \\ &\stackrel{(173)}{=} \bigcup_{i=1}^x \left(\bigcup_{m=1}^z \left(\left(\bigcap_{j=1}^{y_i} \Gamma_{ij} \right) \prec \left(\bigcap_{n=1}^{t_m} \Upsilon_{mn} \right) \right) \right) \\ &\stackrel{(171)}{=} \bigcup_{i=1}^x \left(\bigcup_{m=1}^z \left(\bigcap_{j=1}^{y_i} \left(\Gamma_{ij} \prec \left(\bigcap_{n=1}^{t_m} \Upsilon_{mn} \right) \right) \right) \right) \\ &\stackrel{(171)}{=} \bigcup_{i=1}^x \left(\bigcup_{m=1}^z \left(\bigcap_{j=1}^{y_i} \left(\bigcap_{n=1}^{t_m} \Gamma_{ij} \prec \Upsilon_{mn} \right) \right) \right) \\ &= \bigcup_{(i=1, m=1)}^{(x, z)} \left(\bigcap_{(j=1, n=1)}^{(y_i, t_m)} \Gamma_{ij} \prec \Upsilon_{mn} \right) \\ &= \bigcup_{u=1}^{x \times z} \left(\bigcap_{v=1}^{y_i \times t_m} \Theta_{uv} \right). \quad (226) \end{aligned}$$

In the above rewriting, associativity of intersections and unions was used to go to the penultimate line, and then the definition of the Θ_{uv} 's was injected to go the last line. Note that compared to the intersection and union cases, there is no risk of redundancies since the permutations of Γ_{ij} and Υ_{mn} run over different sets ($\{A, \dots, K\}$ and $\{L, \dots, R\}$ respectively). As each Θ_{ab} is a "prec chain", it is direct to see that the last line is a normal form of $\Gamma \prec \Upsilon$.

Using relations (56) (proven in App. A.5.2), the no signaling composition of two normal forms, $\Gamma \otimes \Upsilon$, as well as their transformation of a normal form into a normal form, $\Gamma \rightarrow \Upsilon$ can be expressed in terms of intersections, unions, negations or prec: $\Gamma \otimes \Upsilon = (\Gamma \prec \Upsilon) \cap (\Gamma \succ \Upsilon)$ and $\Gamma \rightarrow \Upsilon = (\bar{\Gamma} \prec \Upsilon) \cap (\bar{\Gamma} \succ \Upsilon)$. Thus, they can be put in normal form as well because of the above discussion.

Therefore, we showed that the negation of a normal form, the intersection and union of two normal forms, as well as the composition of two normal forms using $\otimes$, $\prec$, and $\rightarrow$ can all be rewritten into a normal form. The proof is completed by noticing that a single base projector like $\mathcal{P}_A$ is by definition in a normal form.

D.8 Proof of Theorem 3

Using Lemma 6, the equivalence between (77) and (72) is almost immediate: inject the result on each node, $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-network})} \equiv (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \prec \dots \prec (\mathcal{I}_{A_{2n-2}} \rightarrow \mathcal{I}_{A_{2n-1}}) \stackrel{(76a)}{=} (\bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}) \prec \dots \prec (\bar{\mathcal{I}}_{A_{2n-2}} \prec \mathcal{I}_{A_{2n-1}})$, then use the associativity of the prec.

This in turn can be used to recursively prove the equivalence between (77) and (74). Indeed, observe that the 1-comb is characterized by $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}$ which is equivalent to

$$\mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2\text{-comb})} \equiv \mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} \stackrel{(76a)}{=} \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}. \quad (227)$$

Then, each iteration of combs satisfies the latter condition of Eq. (76a) on the right side of the $\rightarrow$. E.g, for $\mathcal{P}_{\mathcal{A}_{\text{state}}}^{(4\text{-comb})}$:

$$\begin{aligned} \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(4\text{-comb})} &\equiv ((\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \rightarrow \mathcal{I}_{A_2}) \rightarrow \mathcal{I}_{A_3} \\ &= ((\bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}) \rightarrow \mathcal{I}_{A_2}) \rightarrow \mathcal{I}_{A_3} \\ &= (\bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \mathcal{I}_{A_2}) \rightarrow \mathcal{I}_{A_3} \\ &= (\mathcal{I}_{A_0} \prec \bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2}) \rightarrow \mathcal{I}_{A_3} \\ &= \bar{\mathcal{I}}_{A_0} \prec \bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \mathcal{I}_{A_3} \\ &= \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \bar{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_3}. \quad (228) \end{aligned}$$

Where the associativity of one-way signaling composition, Eq. (160), and distributed negation, Eq. (161), has been used to simplify in between each step. The proof for the n -comb directly follows by induction on the above computation: suppose it holds for n , $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-network})} = \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})}$ then for $n+1$ the projector is

$$\begin{aligned} \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2(n+1)\text{-comb})} &= (\mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})} \rightarrow \mathcal{I}_{A_{2n}}) \rightarrow \mathcal{I}_{A_{2n+1}} \\ &= (\bar{\mathcal{P}}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})} \prec \mathcal{I}_{A_{2n}}) \rightarrow \mathcal{I}_{A_{2n+1}} \\ &= \bar{\mathcal{P}}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})} \prec \mathcal{I}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}} \\ &= \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2n\text{-comb})} \prec \bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}} \\ &= \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-network})} \prec (\mathcal{I}_{A_{2n}} \rightarrow \mathcal{I}_{A_{2n+1}}) \\ &\equiv \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{((n+1)\text{-network})}, \quad (229) \end{aligned}$$

where the hypothesis was injected in between the antepenultimate and penultimate lines as well as identity $\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}} = (\mathcal{I}_{A_{2n}} \rightarrow \mathcal{I}_{A_{2n+1}})$.

Next, the equivalence between (77) and (73) is also proven by induction. It holds by definition for the $n=1$ case, suppose it holds for n , $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})} = \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \dots \prec \bar{\mathcal{I}}_{A_{2n-2}} \prec \mathcal{I}_{A_{2n-1}}$, and define the relabelling $\mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'} \equiv \bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \dots \prec \bar{\mathcal{I}}_{A_{2n-1}} \prec \mathcal{I}_{A_{2n}}$ where all indices have been incremented by 1. Then,

$$\begin{aligned} \mathcal{P}_{\mathcal{A}_{\text{state}}}^{(2(n+1)\text{-comb})} &\equiv (\dots (\mathcal{I}_{A_n} \rightarrow \mathcal{I}_{A_{n+1}}) \rightarrow \dots) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_{2n+1}}) \\ &= \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'} \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_{2n+1}}) \\ &\stackrel{(144)}{=} (\mathcal{I}_{A_0} \otimes \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'}) \rightarrow \mathcal{I}_{A_{2n+1}} \\ &\stackrel{(76b)}{=} (\mathcal{I}_{A_0} \prec \mathcal{P}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'}) \rightarrow \mathcal{I}_{A_{2n+1}} \\ &\stackrel{(76a)}{=} \bar{\mathcal{I}}_{A_0} \prec \bar{\mathcal{P}}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'} \prec \mathcal{I}_{A_{2n+1}} \\ &= \bar{\mathcal{I}}_{A_0} \prec \bar{\mathcal{P}}_{\mathcal{A}_{\text{channel}}}^{(\text{n-comb})'} \prec \mathcal{I}_{A_{2n+1}} \\ &= \bar{\mathcal{I}}_{A_0} \prec \bar{\mathcal{I}}_{A_1} \prec \dots \prec \bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}} \\ &= \bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \dots \prec \bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}}. \quad (230) \end{aligned}$$

Here, the *uncurrying rule* (144) (first proven for types [4]; it is derived for projectors in App. A.4.2) was used between the second and third lines as a computational shortcut.

Finally, the equivalence between (77) and (75) follows by induction as well. In the case $n=1$, it is proven by writing explicitly the content of the projector (77),

$$\begin{aligned} (\bar{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}) \{M\} &= M \\ (\mathcal{I}_{A_0} \otimes \mathcal{I}_{A_1} - \mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1} + \mathcal{D}_{A_0} \otimes \mathcal{D}_{A_1}) \{M\} &= M \\ &= M - \frac{\mathbb{1}_{A_1}}{d_{A_1}} \otimes \text{Tr}_{A_1} [M] \\ &\quad + \frac{\mathbb{1}_{A_0}}{d_{A_0}} \otimes \text{Tr}_{A_0} \left[\frac{\mathbb{1}_{A_1}}{d_{A_1}} \otimes \text{Tr}_{A_1} [M] \right] = M \\ \frac{\mathbb{1}_{A_0 A_1}}{d_{A_0 A_1}} \otimes \text{Tr}_{A_0 A_1} [M] &= \frac{\mathbb{1}_{A_1}}{d_{A_1}} \otimes \text{Tr}_{A_1} [M] \\ \text{Tr}_{A_1} [M] &= \frac{1}{d_{A_0}} \text{Tr}_{A_0 A_1} [M] \mathbb{1}_{A_0}. \quad (231) \end{aligned}$$

Suppose it holds for n nodes, i.e., $\mathcal{P}^{(2n)} \{M^{(n)}\} = M^{(n)} \iff \text{Tr}_{A_1} [M^{(1)}] = \frac{1}{d_{A_0}} \text{Tr}_{A_0 A_1} [M^{(1)}] \mathbb{1}_{A_0} \wedge \dots \wedge \text{Tr}_{A_{2n-1}} [M] = \frac{1}{d_{A_{2n-2}}} \text{Tr}_{A_{2n-2} A_{2n-1}} [M] \otimes \mathbb{1}_{A_{2n-1}}$. Then, for

$n + 1$ nodes, let $M^{(n+1)} \equiv M$, and

$$\begin{aligned}
M &= \mathcal{P}^{(2n+2)}\{M\} \\
&= [\mathcal{P}^{(2n)} \prec (\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}})]\{M\} \\
&= [(\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) \otimes (\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}}) \\
&\quad - \bar{\mathcal{P}}^{(2n)} \otimes \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}} \\
&\quad + \mathcal{D}_{A_0} \otimes \mathcal{D}_{A_1} \otimes \dots \otimes \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}}]\{M\} \\
&= [(\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) \otimes (\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}}) \\
&\quad - \mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}} \otimes \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}} + \\
&\quad \mathcal{P}^{(2n)} \otimes \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}}]\{M\}. \quad (232)
\end{aligned}$$

Using this last equality, one can regroup terms as

$$\begin{aligned}
0 &= [(\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) \otimes \\
&\quad ((\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}}) - \mathcal{I}_{A_{2n}} \otimes \mathcal{I}_{A_{2n+1}})]\{M\} \\
&\quad - [((\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) - \mathcal{P}^{(2n)}) \otimes \\
&\quad \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}}]\{M\}. \quad (233)
\end{aligned}$$

This defines two projectors with zero intersection, therefore each piece in square brackets must be zero independently of the other. The first piece, $0 = [(\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) \otimes ((\bar{\mathcal{I}}_{A_{2n}} \prec \mathcal{I}_{A_{2n+1}}) - \mathcal{I}_{A_{2n}} \otimes \mathcal{I}_{A_{2n+1}})]\{M\}$ is exactly equation (231) applied on systems $A_{2n+1}A_{2n}$. Whereas the second piece, $0 = [((\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) - \mathcal{P}^{(2n)}) \otimes \mathcal{D}_{A_{2n}} \otimes \mathcal{D}_{A_{2n+1}}]\{M\}$ can be recast into $0 = [((\mathcal{I}_{A_0} \otimes \dots \otimes \mathcal{I}_{A_{2n-1}}) - \mathcal{P}^{(2n)})\{\text{Tr}_{A_{2n}A_{2n+1}}[M]\}]$. Using $M^{(n)} \equiv \text{Tr}_{A_{2n}A_{2n+1}}[M]$, this last equation must contain by hypothesis the n other causality conditions. Therefore, the $n + 1$ causality conditions have been recovered from the projector, completing the proof.

E Detailed examples for Section 3 – (Some) state structures built with quantum theory as base state structure

Here we present some examples of state structures that can be built using the three operations $\{\bar{\cdot}, \otimes, \rightarrow\}$, respectively called *negation*, *tensor*, and *transformation* and presented in Sec. 3. The base state structures $\mathcal{A}, \mathcal{B}, \dots$ associated with each party in the following examples are assumed

to be the set of quantum states, to which correspond projectors $\mathcal{I}_A, \mathcal{I}_B, \dots$

With these base state structures, we consider first single-partite POVM measurements, which are the resolutions of the unit effect (or functional) $\bar{\mathcal{A}}$ characterized by $\bar{\mathcal{I}}_A$ as an example of negation. Then the bipartite quantum states $\mathcal{A} \otimes \mathcal{B}$, characterized by $\mathcal{I}_A \otimes \mathcal{I}_B$, are considered as an example of tensor, followed by bipartite functional $\bar{\mathcal{A}} \otimes \bar{\mathcal{B}}$, as an example of tensor and negation. Next, we consider quantum channels $\mathcal{A} \rightarrow \mathcal{B}$, characterized by $\mathcal{I}_A \rightarrow \mathcal{I}_B$ as an example of transformation, followed by single-partite process matrix $\bar{\mathcal{A}} \rightarrow \bar{\mathcal{B}}$ as an example of transformation and negation.

We then move on to “higher-order theory”, interpreting the transformations as higher-order states. Our aim is to present the difference between the set of bipartite quantum channels, $(\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$, and its subset of no signaling channels, $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$. The former indeed has the type of a first-order transformation: it transforms the bipartite states $\mathcal{A}_0 \otimes \mathcal{B}_0$ into bipartite states $\mathcal{A}_1 \otimes \mathcal{B}_1$. The latter is also of the first order, but built differently: it is a tensor product of transformations rather than a transformation between tensor products. This difference is used to define a higher-order: the state structure $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$ is a first-order transformation with respect to the state structures of quantum states (characterized by $\mathcal{I}_X$), but it can be interpreted with respect to the state structure of quantum channels (characterized by $\mathcal{I}_{X_0} \rightarrow \mathcal{I}_{X_1}$), effectively moving its interpretation to a higher-order transformation theory. In such a case, the state structure $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$ is the parallel composition of two systems A and B of the base type (i.e. represented by quantum channels); the no signaling channels are seen as bipartite higher-order states. The last example then concerns the functionals normalized on these, which are the bipartite process matrices.

E.1 Single party quantum theory

This is the simplest example of a theory that can be built assuming Proposition 1. In this case, we have a state structure $\mathcal{A}$ whose operator system spans the whole of $\mathcal{L}(\mathcal{H}^A)$, which will constitute the states. To it, we associate its negation $\bar{\mathcal{A}}$,

which constitutes the unit effects (or deterministic functionals) interpreted as ‘the need of a state of $\mathcal{A}$ to obtain a probability of 1’. This results in the state being the regular notion of the quantum state in density matrix form, as defined in Eq. (8),

$$V \geq 0, \quad (234a)$$

$$\text{Tr}[V] = 1, \quad (234b)$$

$$\mathcal{I}\{V\} = V. \quad (234c)$$

To it corresponds the regular notion of effect as $\overline{\mathcal{A}} = \{\mathbb{1}\}$ because Proposition 1 implies that the valid $N \in \overline{\mathcal{A}}$ have to obey

$$N \geq 0, \quad (235a)$$

$$\text{Tr}[N] = d_A, \quad (235b)$$

$$\mathcal{D}\{N\} = N, \quad (235c)$$

where we used $\overline{\mathcal{I}} = \mathcal{I} - \mathcal{I} + \mathcal{D} = \mathcal{D}$. These requirements single out $N = \mathbb{1}$. That is to say, the set of destructive measurements of any quantum state whose outcome can be predicted with certainty (whence the terminology *deterministic* functional) is the singleton $\{\mathbb{1}\}$, i.e. the discard. The destructive measurement whose outcome can be predicted only up to a probability, the effects (or probabilistic functionals), are (represented by) the operators resolving $\mathbb{1}$; we have derived the POVM formalism. By construction, the probability rule in that case indeed recovers the Born rule (3):

$$p(i|V, N) = \langle N_i, V \rangle, \quad (236)$$

where each N_i belongs to a family $\{N_i\}$ resolving $\mathbb{1}$. With respect to the introductory example of Sec. 3, only the notation has been changed: $(\rho, \mathbb{1}, E_i) \mapsto (V, N, N_i)$.

E.2 Bipartite quantum theory

Using Definition 5, the bipartite quantum theory is an example of a theory featuring a tensor of projectors. From the previous example of single-party quantum theory, we can define a bipartite theory on space $\mathcal{L}(\mathcal{H}^X)$ by requiring a bipartition $\mathcal{L}(\mathcal{H}^X) \cong \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ so that there exist states upon which bipartite measurements are well-defined (see the discussion around Corollary 2 as well as App. C). The one-party example has shown that local measurements of, say, Alice

resolve the state structure $\overline{\mathcal{A}}$ as in Eqs. (234). Assume Bob has a similarly defined $\mathcal{B}$ so that the parallel composition of the measurements of Alice and Bob are resolving $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$. By Definition 5, this set is characterized by

$$M \geq 0, \quad (237a)$$

$$\text{Tr}_{AB}[M] = 1, \quad (237b)$$

$$(\mathcal{D}_A \otimes \mathcal{D}_B)\{M\} = M. \quad (237c)$$

It is not hard to see from Eq. (128a) that also in the bipartite case, the effects of quantum theory resolve a single element, $\{M = \mathbb{1} = \mathbb{1}_A \otimes \mathbb{1}_B\}$. Note, however, that the effects can now in general be entangled in the sense that there are resolutions $\{M_i\}$ for which there is no possibility to find a decomposition $\{M_i = \sum_i q_i E_i^A \otimes F_i^B\}$ where, for all q_i and i , $q_i \geq 0$, $\sum_i q_i = 1$ and with E_i (respectively, F_i) being a valid effect resolving an element of $\overline{\mathcal{A}}$ (resp., $\overline{\mathcal{B}}$). A Bell measurement is an instance of such a collection of entangled effects resolving $\mathbb{1}$ in the $d_A = d_B = 2$ case.

We now use Proposition 1 to characterize the valid states. An operator $W \in \mathcal{L}(\mathcal{H}^A \otimes \mathcal{H}^B)$ is a valid state if it belongs to the set $\overline{\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}}$, i.e.

$$W \geq 0, \quad (238a)$$

$$\text{Tr}_{AB}[W] = 1, \quad (238b)$$

$$(\mathcal{D}_A \otimes \mathcal{D}_B)\{W\} = W. \quad (238c)$$

Remark that property (139) actually applies to the projective constraint, so that it can be simplified into

$$(\mathcal{I}_A \otimes \mathcal{I}_B)\{W\} = W. \quad (239)$$

Meaning that $\overline{\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}} \cong \overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$ actually holds for quantum theory¹⁶, and surprisingly it does so only in this case (this is a consequence of Lemma 6, proven in App. B). In other words, in quantum theory the set of valid states normalized on local measurements is exactly the set of no signaling composite states: we recover the intuition that it is impossible for parties measuring a part of a shared quantum state to signal to the other one. What is more surprising is that it is actually the only theory having this property: in general a

¹⁶In categorical language, this is the fact that the category of quantum theory is compact closed, whereas the more general category of higher-order quantum transformations is *-autonomous, see Ref. [9].

state normalized on a pair of local measurements can be used for signaling; this is the physical content of the inclusion $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}} \supseteq \mathcal{A} \otimes \mathcal{B}$ which follows from Eq. (33). This ability to signal is made apparent in the bipartite process matrix and in the biased quantum theory examples below (respectively presented in App. E.5 and F.1).

E.3 Single party quantum channel and process matrix formalisms

The first example of a state structure built using the transformation connector ‘ $\rightarrow$ ’ described in Proposition 2 is the one of quantum channels. With the associated quantum instrument formalism and the single-partite process matrices, they form a higher-order state and effect pair. This is the same pair as the one considered in the introductory example of Sec. 3, rephrased using the language developed in that section (it will be rephrased again in App. G.1 using all of the language developed in this paper). The construction is schematically represented in Fig. 11.

Let $\mathcal{A}_0$ and $\mathcal{A}_1$ be state structures characterized by similar constraints as Eqs. (234). In accordance with Definition 6, we introduce dynamics from $\mathcal{A}_0$ to $\mathcal{A}_1$ as the set of all CPTP maps $\{\mathcal{M}\}$ from $\mathcal{A}_0$ to $\mathcal{A}_1$; this is represented as step 1 in Fig. 11. Call $\mathcal{A}_0 \rightarrow \mathcal{A}_1$ the set of all operators $M \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$ which are the CJ representation of an admissible transformation; these are the M such that for all quantum state $\rho_{A_0} \in \mathcal{A}_0 \subset \mathcal{L}(\mathcal{H}^{A_0})$, there exists a state $\sigma_{A_1} \in \mathcal{A}_1 \subset \mathcal{L}(\mathcal{H}^{A_1})$ so that

$$[\text{Tr}_{A_0} [M \cdot (\rho_{A_0} \otimes \mathbb{1})]]^T = \sigma_{A_1} . \quad (240)$$

By definition, the operators M ’s are (representing) quantum channels, so by Proposition 2, they must satisfy

$$M \geq 0 , \quad (241a)$$

$$\text{Tr} [M] = d_{A_0} , \quad (241b)$$

$$(\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \{M\} = M , \quad (241c)$$

to be valid; this is depicted as step 2 in Fig. 11. To see that these conditions indeed single out the quantum channels in CJ form, remark that the last two conditions can be combined in the more common quantum ‘1-comb’ condition, $\text{Tr}_{A_1} [M] = \mathbb{1}_{A_0}$ (see Ref. [3, 11]; see also the discussion in Sec. 6.2). It is obtained by explicitly

writing the projector

$$\begin{aligned} (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \{M\} &= \overline{\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}} \{M\} \\ &= (\mathcal{I}_{A_0} \otimes \mathcal{I}_{A_1} - \mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1} + \mathcal{D}_{A_0} \otimes \mathcal{D}_{A_1}) \{M\} , \end{aligned} \quad (242)$$

and noticing that

$$(\mathcal{D}_{A_0} \otimes \mathcal{D}_{A_1}) \{M\} = \frac{\text{Tr}_{A_0 A_1} [M]}{d_{A_0} d_{A_1}} \mathbb{1}_{A_0 A_1} . \quad (243)$$

Seeing these as the effects of a higher-order theory, the corresponding set of states is obtained through Prop. 1 (notice we are actually working the proof of Prop. 2 backward; see App. D.2):

$$W \geq 0 , \quad (244a)$$

$$\text{Tr} [W] = d_{A_1} , \quad (244b)$$

$$\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}} \{W\} = W . \quad (244c)$$

This is depicted as step 3 in Fig. 11; this, by definition, is the set of single-partite process matrices, as it is the set of functionals normalized on quantum instruments. Indeed, any resolution $\{M_i\}$ of a channel leads to a probability rule of the form of Eq. (5):

$$p(i|W, M) = \langle M_i , W \rangle . \quad (245)$$

The set of valid W ’s is characterized by projector $\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}} = \overline{\mathcal{I}_{A_0} \otimes \mathcal{I}_{A_1}} = \mathcal{I}_{A_0} \otimes \overline{\mathcal{I}_{A_1}}$. Thus it corresponds to state structure $\mathcal{A}_0 \otimes \overline{\mathcal{A}_1}$, characterized by $\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}$. By Definition 2, it is supposedly in the affine span of the objects of the form $\rho_{A_0} \otimes \mathbb{1}_{A_1}$. However, the remarkable thing here is that, since there is only one element in $\overline{\mathcal{A}_1} = \{\mathbb{1}_{A_1}\}$, there is no need to take the affine span as $\mathcal{A}_0$ is convex closed: in Fig. 11 it means that the pure tensor products of the pink and green objects are exactly the yellow one. In other words, this recovers the known result that a single-partite process matrix reduces to inputting a quantum state and tracing out the output (see the supplementary material of Ref. [8], alternatively see Ref. [34]).

Remark: the no signaling subset of a quantum channel. From the above considerations, it can be guessed that the set of elements of the form $M = \mathbb{1}_{A_0} \otimes \sigma_{A_1}^T$, which is $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1$, is a subset of the valid quantum channels. And indeed they are the ‘trace-and-replace’ channels such that any

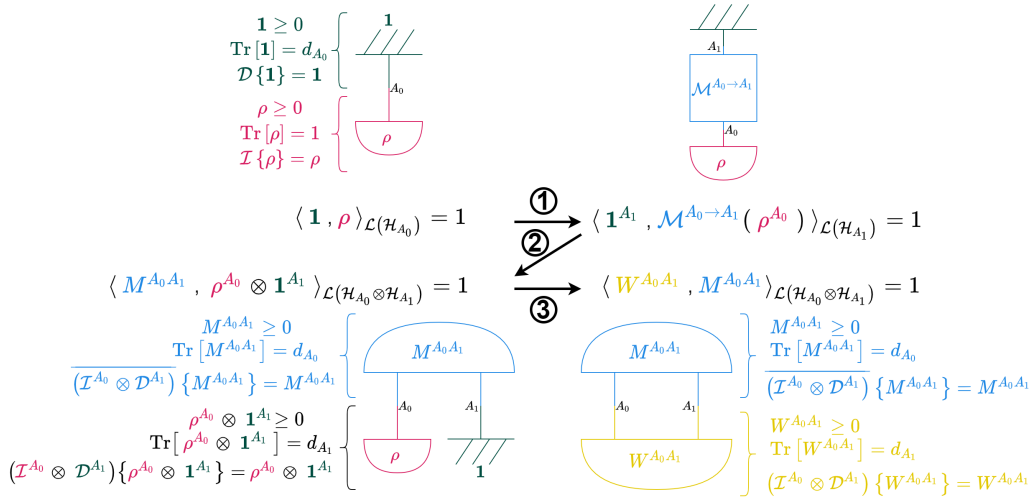


Figure 11: Dynamics-like construction from quantum states to single-partite PM: 1) adding CPTP dynamics to states (in pink) yields channels (in blue), 2) going to the CJ picture and characterizing it using Prop. 2, 3) taking the ‘negation’ of the channels to ‘complete’ the states and effects pair (pink and green) into a single-partite PM (in yellow)

input state ρ_{A_0} is traced out and replaced by σ_{A_1} :

$$\left(\text{Tr}_{A_0} \left[\left(\mathbb{1}_{A_0} \otimes \sigma_{A_1}^T \right) \cdot (\rho_{A_0} \otimes \mathbb{1}_{A_1}) \right] \right)^T = \text{Tr}_{A_0} [\rho_{A_0}] \left(\sigma_{A_1}^T \right)^T = \sigma_{A_1}. \quad (246)$$

In terms of state structures this is the inclusion $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1 \subseteq \mathcal{A}_0 \rightarrow \mathcal{A}_1 = \overline{\mathcal{A}_0} \otimes \overline{\mathcal{A}_1}$ so it is again Eq. (33) at play: the transformation can be seen as another composition of $\overline{\mathcal{A}_0}$ with $\mathcal{A}_1$ that is larger than the one provided by their tensor product, $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1$. But here, compared to the bipartite quantum theory, the inclusion is strict, $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1 \subsetneq \mathcal{A}_0 \rightarrow \mathcal{A}_1 = \overline{\mathcal{A}_0} \otimes \overline{\mathcal{A}_1}$; the ‘trace-and-replace’ are indeed but a specific kind of channels.

Actually, a special property of the subset $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1$ is that neither side of the channel can have a deterministic influence over the other. To see it, treat $\overline{\mathcal{A}_0}$ and $\mathcal{A}_1$ as if they were under the control of different parties so the ‘trace’ part is controlled by A_0 and the ‘replace’ part by A_1 . The resolutions $\{M_i\}$ of M now depend on 2 classical variables: $i \mapsto (j, k)$ such that j is known only by A_0 and k only by A_1 . Any such ‘bipartite’ resolution $\{M_{j,k}\}$ has well-defined ‘single-partite marginals’:

$$\begin{aligned} \sum_j M_{j,k} &= \mathbb{1} \otimes \sigma_k, \\ \sum_k M_{j,k} &= E_j \otimes \sigma^T, \\ \sum_{j,k} M_{j,k} &= \mathbb{1} \otimes \sigma^T, \end{aligned} \quad (247)$$

in the sense that $\{E_j\}$ is a resolution of $\mathbb{1}$ and $\{\sigma_k\}$ is one of σ^T . Indeed, if such partial sums were not possible to obtain, the interpretation as local parties acting probabilistically would fail; one could not interpret $\sum_j M_{j,k}$ (respectively, $\sum_k M_{j,k}$) as a resolution where the action of A_0 (resp., A_1) on her side of the shared channel is deterministic. Such ‘bipartite’ resolution $M_{j,k}$ of a ‘trace-and-replace’ map has interpretation as a ‘measure-and-reprepare’ scenario; the variable j is associated with the measurement outcome seen by A_0 whereas k conditions the choice of reparation made by A_1 .

The probability rule becomes $p(j, k|W, M) = \langle M_{j,k}, W \rangle$ and can be simplified using the fact that W factorises into $\rho \otimes \mathbb{1}$,

$$p(j, k|W, M) = \text{Tr} [M_{j,k} (\rho \otimes \mathbb{1})]. \quad (248)$$

Talking about one side’s ‘deterministic influence’ over the other requires summing over her classical variable. The probabilistic interpretation of such summation is precisely a conditional distribution. For example, after A_0 has acted on her side the ‘reduced’ distribution as seen on A_1 side is $\sum_j p(j, k|W, M) = p(k|j; W, M)$; it is conditioned by ‘quantum variables’ W, M (as well as how M has been resolved into $\{M_{j,k}\}$) and by the actual value of the classical variable j that A_0 has observed. However, for an $M \in \overline{\mathcal{A}_0} \otimes \mathcal{A}_1 \subset \mathcal{A}_0 \rightarrow \mathcal{A}_1$, the conditional distributions seen by each side drastically simplify

because of Eqs. (247),

$$\begin{aligned} \sum_j p(j, k|W, M) &= \text{Tr} [\sigma_k] ; \\ \sum_k p(j, k|W, M) &= \text{Tr} [E_j \rho] . \end{aligned} \quad (249)$$

Consequently, the classical variables are independent of each other, $p(k|j; W, M) = p(k|W, M)$ and vice-versa. In other words, the knowledge of j cannot help determine k , and vice-versa. In such a channel no information can be deterministically passed from the input side at A_0 to the output side at A_1 ; the subset $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1$ is dubbed the ‘no signaling’ subset of $\mathcal{A}_0 \rightarrow \mathcal{A}_1$ for this reason. This notion of ‘no signaling’ as the independence of conditional probabilities is generalized to any state structure in Sec. 4.2.2 under Definition 7, and it is generally proven to correspond to the tensor composition of state structures in Proposition 5.

Remark: channels are no signaling from output to input. Another occurrence of no signaling in state structures is the property of quantum channels to be automatically no signaling from the output to the input. So far we considered M to represent the action of the parties and W to be the environment, in the sense that M was being resolved and W was deterministic and fixed. Reversing the picture, the transformation M in $\overline{\mathcal{A}_0} \otimes \mathcal{A}_1$ can be seen as a channel shared by two parties A_0 and A_1 trying to communicate through it. The description of their joint intervention is thus an element of $\mathcal{A}_0 \otimes \overline{\mathcal{A}_1}$ (or a resolution of one such element if they are acting probabilistically). As in the above remark, the probabilistic dependency of W can be split between two variables j and k , each known only by one of the parties. A_0 prepares a state according to j , inputs it through the channel, and then A_1 measures its output. In that case, however, party A_1 ’s action is to perform a POVM, but it means she can apply only one deterministic (or unit) effect at the output of a quantum channel (i.e. $\mathbb{1}_{A_1}$). It is impossible for her to deterministically influence the input of the channel by her action on the output as there are no two different deterministic effects N and N' such that $\text{Tr}_{A_1} [M \cdot (\mathbb{1} \otimes N)] \neq \text{Tr}_{A_1} [M \cdot (\mathbb{1} \otimes N')]$; no matter what A_1 does deterministically, A_0 always sees the same thing on her end of the channel. This absence of deterministic influence is

then to be interpreted as the no signaling from the output to the input condition (58a) defined in App. C. This property is not true for transformations in general (again, it follows as a consequence of Lemma 6): if M was not a quantum channel but, say, a transformation between process matrices, the party at the output could have deterministic influence over what the input sees. This observation is a motivation for the introduction of the one-way signaling composition of two state structures in Sec. 5 as well as the discussion about why the quantum comb formalism is “accidentally” one-way signaling in Sec. 6.2.

E.4 Bipartite quantum channel theory and the no signaling channels

This example is one of the simplest that mixes the tensor and the transformation, i.e. Definition 5 and Proposition 2. A bipartite quantum channel is a CPTP map $\mathcal{M}$ between bipartite quantum states. For $\mathcal{M} : \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{B_0}) \rightarrow \mathcal{L}(\mathcal{H}^{A_1} \otimes \mathcal{H}^{B_1})$, the set of all such channels corresponds to the state structure $(\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$. We will see in this example that the ordering of parentheses and connectors $\otimes, \rightarrow$ is important, and that this semantic difference underlies the difference between the bipartite channels and their no signaling subset. Like the previous subsection, this example also hints at how the notion of no signaling is generalized in Sec. 5 and especially how it relates to the notion of no signaling at the level of the outcomes probability distribution that is developed in App. C.

Bipartite quantum channels A bipartite channel is example E.3 in the bipartite case, refer to the top right of Fig. 12 for an illustration. The deterministic probability rule reads:

$$\langle \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_1}, \mathcal{M}(\rho_{A_0 B_0}) \rangle = 1, \quad (250)$$

which holds for $\mathbb{1}_{A_1} \otimes \mathbb{1}_{B_1} \in \overline{\mathcal{A}_1} \otimes \mathcal{B}_1$ and for all $\rho_{A_0 B_0} \in \mathcal{A}_0 \otimes \mathcal{B}_0$. In the CJ picture, this is

$$\langle M, \rho_{A_0 B_0} \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_1} \rangle = 1, \quad (251)$$

with $M \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{B_0} \otimes \mathcal{H}^{A_1} \otimes \mathcal{H}^{B_1})$, leading to the characterization of valid M 's as:

$$M \geq 0, \quad (252a)$$

$$\text{Tr}[M] = d_{A_0} d_{B_0}, \quad (252b)$$

$$((\mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0}) \rightarrow (\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1})) \{M\} = M. \quad (252c)$$

The projector is $(\mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0}) \rightarrow (\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1}) = \mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0} \otimes \mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1} - \mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0} \otimes \overline{\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1}} + \mathcal{D}_{A_0} \otimes \mathcal{D}_{B_0} \otimes \mathcal{D}_{A_1} \otimes \mathcal{D}_{B_1}$; It can be further simplified by noticing that $\overline{\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1}} = \mathcal{D}_{A_1} \otimes \mathcal{D}_{B_1}$. This simplification then gives condition $M - \text{Tr}_{A_1 B_1}[M] \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_1} + \mathbb{1}_{A_0} \otimes \mathbb{1}_{B_0} \otimes \mathbb{1}_{A_1} \otimes \mathbb{1}_{B_1} = M$ which, once the trace is fixed, is equivalent to the usual condition $\text{Tr}_{A_1 B_1}[M] = \mathbb{1}_{A_0} \otimes \mathbb{1}_{B_0}$. One indeed recovers the same condition as for the single-partite case but applied on two subsystems.

Remark: Bipartite channels from the composition of two channels. We defined the state structure of bipartite channels as the set of transformations between bipartite states, one can nonetheless wonder if it can be defined as the composition of two state structures of single-partite channels instead. In other words, can we define the state structure $(\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$ by combining state structures $\mathcal{A}_0 \rightarrow \mathcal{A}_1$ and $\mathcal{B}_0 \rightarrow \mathcal{B}_1$? The tensor product appears as a good candidate to do so. However, it will be shown next that it only yields the subset of no signaling bipartite quantum channels. As it is discussed in Sec. 4.2.2 under Eq. (33), there exists a larger composition than the tensor, based on the non-commutation of the negation with the tensor. It amounts to taking the negation of the tensor product of the negations or, equivalently, to taking the transformation with a negated input. In symbols, this last sentence amounts to the following inclusion:

$$\mathcal{A} \otimes \mathcal{B} \subseteq \overline{\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}} = \overline{\mathcal{A}} \rightarrow \mathcal{B} = \mathcal{A} \leftarrow \overline{\mathcal{B}}. \quad (253)$$

So is the right-hand side the proper composition? It turns out that it is for $\mathcal{A} = \mathcal{A}_0 \rightarrow \mathcal{A}_1$ and $\mathcal{B} = \mathcal{B}_0 \rightarrow \mathcal{B}_1$ being the state structures of quantum channels. This can be shown by some projector algebra and, as for the case of bipartite quantum states above, this crucially relies on the fact that the identity projector on two systems is the same as the tensor of two identities as in Eq. (128b). The projector associated to state

structure $\overline{\overline{\mathcal{A}_0 \rightarrow \mathcal{A}_1} \otimes \overline{\mathcal{B}_0 \rightarrow \mathcal{B}_1}}$ can undergo the following rewriting:

$$\begin{aligned} \overline{\overline{\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}} \otimes \overline{\mathcal{I}_{B_0} \rightarrow \mathcal{I}_{B_1}}} \\ &= \overline{(\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}) \otimes (\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1})} \\ &= \overline{\mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0} \otimes \mathcal{D}_{A_1} \otimes \mathcal{D}_{B_1}} \\ &\stackrel{(128a)}{=} \overline{\mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0} \otimes \overline{\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1}}} \\ &= (\mathcal{I}_{A_0} \otimes \mathcal{I}_{B_0}) \rightarrow (\mathcal{I}_{A_1} \otimes \mathcal{I}_{B_1}), \end{aligned} \quad (254)$$

and this last line is indeed the projector associated with bipartite quantum channels.

The fact that the bipartite quantum channels are equivalent to the composition of single-partite channels is not trivial. This only happens in a handful of cases in which the base state structures are quantum, i.e. when the projectors associated with subsystems are the identity or depolarizing. (The cases are inferred from Eq. (139) which is again a consequence of Lemma 6 presented in Sec. (6.2).) In general, the following inclusion is tight,

$$\begin{aligned} (\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1) \\ \subseteq (\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1) \\ \subseteq \overline{\overline{(\mathcal{A}_0 \rightarrow \mathcal{A}_1)} \otimes \overline{(\mathcal{B}_0 \rightarrow \mathcal{B}_1)}}, \end{aligned} \quad (255)$$

so a bipartite transformation is a genuinely different state structure than a composition of transformations.

No signaling bipartite channels. A concrete example of the above inclusions is obtained by comparing (the state structures of) the tensor product of quantum channels to a bipartite quantum channel, i.e. $(\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$ compared to $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$ in the case of base state structures of quantum states. The first case, as we have shown, corresponds to the set of bipartite channels, whereas the second is the set of no signaling [18] (also known as causal [17]) channels as we will now show. The tightness of the inclusion can be proven with a dash of the algebra of projectors: using Eq. (255), the bipartite channels are associated with projector $\overline{(\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}) \otimes (\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1})}$ whereas $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$ is associated to $\overline{\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}} \otimes \overline{\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1}}$ so Eqs. (134) and (139) can be used:

$$\begin{aligned} \overline{(\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}) \otimes (\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1})} \\ \supset \overline{\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}} \otimes \overline{\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1}}. \end{aligned} \quad (256)$$

Next, we prove that it is indeed the set of no signaling channels, that is, the subset of channels forbidding deterministic signaling from Alice's side to Bob's side and vice-versa. We already mentioned in Remark E.3 that channels are no signaling from the output to input and that local operation on a channel is a single-partite PM, corresponding to a state preparation followed by a discard. Putting these two pieces of information together, what is to be shown here is that the choice of input state on Alice's side cannot induce a deterministic influence on Bob's measurement and vice-versa. In equation, Alice and Bob share a bipartite channel M , and each acts locally by state preparation and measurement. Call x and y the settings of respectively Alice and Bob, and a and b their outcomes. That is, depending on some input x , Alice prepares some state $\rho_{|x}$ that she inputs in her side of the channel, A_0 . Then, she measures at the output A_1 some POVM $\{E_{a|x}\}$ which choice can also depend on x , and sees outcome a . Bob does the same at B_0 with state $\sigma_{|y}$ at B_0 and POVM $\{F_{b|y}\}$ at B_1 . The probability rule then reads

$$p(a, b|x, y) = \text{Tr} \left[M \cdot \left(\rho_{|x} \otimes \sigma_{|y} \otimes E_{a|x}^T \otimes F_{b|y}^T \right) \right]. \quad (257)$$

(We omit reference to the 'quantum variables' M , ρ , σ , etc. in the probability distribution to lessen clutter, as we are only interested in the correlations between the classical variables a, b, x , and y .) Alice is no signaling to Bob if her choice of setting x cannot deterministically influence his measurement outcome, that is (see e.g., [18])

$$\forall x, x', \sum_a p(a, b|x, y) = \sum_a p(a, b|x', y). \quad (258)$$

The same way Bob is no signaling if

$$\forall y, y', \sum_b p(a, b|x, y) = \sum_b p(a, b|x, y'). \quad (259)$$

In terms of the channel, these two conditions can be shown to be equivalent to (see [3]):

$$\text{Tr}_{A_1} [M] = \mathbb{1}_{A_0} \otimes \text{Tr}_{A_0 A_1} [M]; \quad (260a)$$

$$\text{Tr}_{B_1} [M] = \mathbb{1}_{B_0} \otimes \text{Tr}_{B_0 B_1} [M]. \quad (260b)$$

which form a stronger constraint than the channel condition $\text{Tr}_{A_1 B_1} [M] = \mathbb{1}_{A_0} \otimes \mathbb{1}_{B_0}$.

These two conditions are exactly those encoded in the projector of the state structure $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$, which reads

$(\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \otimes (\mathcal{I}_{B_0} \rightarrow \mathcal{I}_{B_1}) = \overline{\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}} \otimes \overline{\mathcal{I}_{B_0} \otimes \mathcal{D}_{B_1}}$. Indeed, notice that this is the composition of two projectors acting on different subspaces (and thus commuting):

$$\begin{aligned} (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \otimes (\mathcal{I}_{B_0} \rightarrow \mathcal{I}_{B_1}) &= \\ ((\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \otimes \mathcal{I}_{B_0} \otimes \mathcal{I}_{B_1}) \cap \\ (\mathcal{I}_{A_0} \otimes \mathcal{I}_{A_1} \otimes (\mathcal{I}_{B_0} \rightarrow \mathcal{I}_{B_1})) &, \end{aligned} \quad (261)$$

Where the $\cap$ symbol represents a composition of commuting projectors as defined in Sec. 4 (see App. A.2 for a review of its properties). Hence, each of them enforces a condition independent of the other. The first one, $(\overline{\mathcal{I}_{A_0} \otimes \mathcal{D}_{A_1}} \otimes \mathcal{I}_{B_0} \otimes \mathcal{I}_{B_1}) \{M\} = M$, is explicitly

$$M - \mathbb{1}_{A_1} \otimes \text{Tr}_{A_1} [M] + \mathbb{1}_{A_0} \otimes \mathbb{1}_{A_1} \otimes \text{Tr}_{A_0 A_1} [M] = M, \quad (262)$$

which is equivalent to condition (260a), and thus to (258). The same way, $(\mathcal{I}_{A_0} \otimes \mathcal{I}_{A_1} \otimes \overline{\mathcal{I}_{B_0} \rightarrow \mathcal{D}_{B_1}})$ can be shown equivalent to condition (260b) and thus to (259).

This example of no signaling bipartite channel illustrates how Definition 5 is exactly a composition of two state structures done in a manner that forbids a deterministic influence of each state structure over the other. In the CJ picture, this definition tells us that the set of these channels obeys

$$M \geq 0; \quad (263a)$$

$$\text{Tr} [M] = d_{A_0} d_{B_0}; \quad (263b)$$

$$((\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}) \otimes (\mathcal{I}_{B_0} \rightarrow \mathcal{I}_{B_1})) \{M\} = M; \quad (263c)$$

or, equivalently, is the set of trace-normalized positive operators in the affine hull of tensor products of single-partite channels, i.e.,

$$M = \sum_i q_i M_i^A \otimes M_i^B, \quad (264)$$

where $M \geq 0$, $q_i \in \mathbb{R}$, $\sum_i q_i = 1$, each $M_i^A \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$ obeys single quantum channel conditions (241), and so does each $M_i^B \in \mathcal{L}(\mathcal{H}^{B_0} \otimes \mathcal{H}^{B_1})$. Generalizing this notion of no signaling to general state structures is the purpose of Sec. 5 in which the tensor product of state structure is effectively proven to be no signaling composition, see Corollary 2.

Remark: separability and localizability.

Notice that the projective methods are linear constraints, so while they can give information about the deterministic signaling structure, nothing can be said about separability or localizability. In the context of a no signaling quantum channel, separability will be the analog of separable quantum states: the subset of channels whose decomposition like Eq. (264) is a convex sum, for which each $q_i \in [0, 1]$. Localizability on the other hand is the possibility of obtaining the channel as local operations applied on a joined ancillary entangled state [17]. It is known that separable and localizable channels are subsets of the set of no signaling channels. Yet, the characterization of these sets is much more involved than working out the subspaces they span, as their constraints are non-linear in the operators. Generalizing separability and localizability to higher-order quantum transformations is left open as a future research direction.

E.5 Bipartite no signaling channel and process matrix formalism.

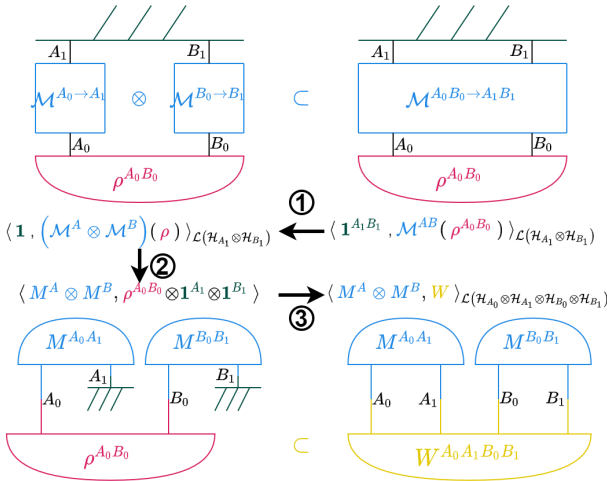


Figure 12: Gleason-like construction of a bipartite PM: starting from a bipartite channel normalized on a state (in blue and pink, top right), 1) Restrict the set of channels to local operations of Alice and Bob subset (in blue, top left); 2) Go to the CJ picture so that the channels are interpreted as two effects composed in tensor (in blue, bottom left); 3) Extend the allowed higher-order ‘state’ to the most general operator normalized on a tensor product of channels, the bipartite PM (in yellow, bottom right)

This last example mixes all three operations on state structures (i.e., $\otimes$, $\rightarrow$, and $\bar{\cdot}$) in order to

recover the bipartite process matrix framework from the four base state structures of quantum states. With respect to the characterization of no signaling channels in the previous section, the bipartite process matrices are nothing short of the functionals normalized on this set [8, 11]. We will do this construction step by step, relating all the other structures characterized in this appendix.

The construction is similar to that of Sec. E.3 but assumes the induced dynamics to be bipartite to start with. Refer to Fig. 12: Starting with the regular bipartite quantum theory as in Sec. E.2 (top right corner, with states in pink and effects in green), one first considers no signaling at the level of effects to argue for bipartite states, $\rho \in \mathcal{A}_0 \otimes \mathcal{B}_0$ (as we have seen, this is automatically the case in the case of quantum theory). Then, the theory is extended by allowing quantum channels (top right corner, the set of channels $\mathcal{M}$ is in blue) as in Sec. E.4 so that these maps correspond to the state structure $(\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$.

Step 1 (going to the top left situation of Fig. 12) consists of restricting the set of valid bipartite channels to the set of no signaling channels, so as to guarantee that parties A and B are causally disconnected. In other words, one considers another requirement of no signaling, at the level of the transformation this time. The heuristic of which is straightforward: if there are local effects, the dynamics may be constrained to local operations as well because e.g. Alice and Bob’s labs are space-like separated. Another way to put it: Alice and Bob are each allowed to do any physical transformation on their qudit before measuring, but no global transformation is allowed. Both points of view conclude that the set of channels is restricted to its local subset spanned by the tensor product of local transformations. Step 2 (going to the bottom left situation of Fig. 12) is to pass to the CJ picture so to obtain their characterization. The operations of Alice and Bob have the form $M^{A_0 A_1} \otimes M^{B_0 B_1} \in (\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1) \subset (\mathcal{A}_0 \otimes \mathcal{B}_0) \rightarrow (\mathcal{A}_1 \otimes \mathcal{B}_1)$ (or, more generally, the affine span of these pure tensors, see Sec. E.4). Finally, step 3 (going to the bottom right situation of Fig. 12) is to extend the set of states and effects to all functionals as in Proposition 1. By doing so, we have inferred from a few heuristics what is the set of bipartite process matrices $\{W\}$, and that its state structure is $(\mathcal{A}_0 \rightarrow \mathcal{A}_1) \otimes (\mathcal{B}_0 \rightarrow \mathcal{B}_1)$, which characteriza-

tion directly ensues (this is done explicitly in Ref. [11]).

Abstractly, this way of building a higher-order theory consists in inducing a no signaling composition of state structures (a tensor product), then normalizing a functional on it (a negation). We will nickname this kind of construction a ‘Gleason-like’ construction as opposed to a ‘dynamics-like’ construction that involves defining dynamics of dynamics (nested transformations; i.e., supermaps). The dynamics-like way of defining higher-order transformations is considered in Sec. G. This choice of nomenclature is made for distinguishing between two lines of works that resulted in the same kind of higher-order formalisms featuring ICO, respectively Refs. [8, 10, 21] and Refs. [2, 3, 6, 4, 5]. In a sense, the Gleason-like way is the special case of the dynamics-like way obtained when the output is the trivial system.

F An example with a different base projector – Biased Quantum Theory

One novelty of our approach compared to the type theory is that the base types –generalized into base state structures– do not have to be the set of quantum states. For example, in a state structure like $\mathcal{A} \rightarrow \mathcal{B}$ the base state structures $\mathcal{A}$ and $\mathcal{B}$ can be any state structure. To compare the resulting process theories, one can consider various sets of possible base state structures for a fixed number of subsystems associated with a fixed number of parties. Say Alice’s system is known to be bipartite and quantum, then her possible base state structures can be the elementary one, $\mathcal{A} = \mathcal{A}_0 \otimes \mathcal{A}_1$, or a transformation $\mathcal{A} = \mathcal{A}_0 \rightarrow \mathcal{A}_1$. The same reasoning applies to Bob.

Nonetheless, the set of base state structures in consideration does not actually have to be restricted to only those that can be built from the quantum states (corresponding to the identity projector) and the operations in the algebra. The only constraint implied by the projective characterization is that the set of base projectors associated with a single party has to contain the depolarizing and the identity projectors and that every projector in the set should commute with the other ones. So what about projectors that cannot be built from the algebra but that com-

mute with the other elements?

As mentioned below Proposition 6 in the main text, a physically motivated example is the set of diagonal operators with respect to a fixed basis. These correspond to the dephased states obtained after a measurement in this basis, hence to classical systems to an extent. The projector Δ to this subspace can be checked to commute with $\mathcal{I}$ and $\mathcal{D}$, so that $\mathcal{I} \supset \Delta \supset \mathcal{D}$, and to obey the other properties of a projector on an operator system (84). The set of base state structures for a single-partite system in this case is characterized by $\{\Delta, \overline{\Delta}\}$ instead of $\{\mathcal{I}, \mathcal{D}\}$. Studying higher-order transformations mixing base state structures of classical and quantum theory is left open as a future research direction.

The purpose of this appendix is instead to present another, less physically relevant, projector that cannot be built from $\mathcal{I}$ and the operations of the algebra, so as to study the properties of the multipartite state structures built out of it. This toy model has different signaling properties than transformations built out of the quantum states as we will show, providing some motivation and insights for Sec. 5.

For simplicity let us take a 2-dimensional underlying Hilbert space $\mathcal{H}^A$. Let $\mathcal{P}$ be the projector that restricts the basis elements to $\{\mathbb{1}, \sigma_x, \sigma_y\}$ in the Pauli basis so that we can define $\mathcal{A} = \text{Span}\{\mathbb{1}, \sigma_x, \sigma_y\}$ associated with projector $\mathcal{P}_A$ and, by Proposition 1, $\overline{\mathcal{A}} = \text{Span}\{\mathbb{1}, \sigma_z\}$ is associated with $\overline{\mathcal{P}}_A$. Compare it to quantum theory in 2 dimensions: $\mathcal{A}_{quant.} = \text{Span}\{\mathbb{1}, \sigma_x, \sigma_y, \sigma_z\}$ associated with $\mathcal{I}_A$ and $\overline{\mathcal{A}}_{quant.} = \{\mathbb{1}\}$ with $\mathcal{D}_A$, it is direct to check that $\mathcal{P}_A$ is a projector on operator system and that $\mathcal{I}_A \subset \mathcal{P}_A \subset \mathcal{D}_A$.

The theory characterized by $\mathcal{P}_A$ is a *biased* quantum theory in the sense that the measurement can be ‘biased’ towards an arbitrary element of $\overline{\mathcal{A}}$. For an element $V \in \mathcal{A}$ and one $N \in \overline{\mathcal{A}}$, the probability rule reads

$$p(i|V, N) = \langle N_i, V \rangle, \quad (265)$$

and it can be interpreted in the picture of regular quantum mechanics as a regular measurement followed by a projection onto the state N .

Thus, some deterministic postselection of the measurement in a given basis –here in the computational basis– is allowed in this kind of theory. Nevertheless, the theory by itself does not allow for the usual counter-logical behaviors encountered in theories with postselection because it is

inherently constraining the allowed states into a basis that is quasi-orthogonal to the one of the postselection. Another way to picture it is that the theory allows for a postselection but in a basis that is by construction mutually unbiased with respect to the basis of the state: in the example, Alice can in general postselect in any state of the form $\mathbb{1} + p \sigma_z$, where p real and $p^2 \leq 1$ because of positivity. That is, she can choose to project the state into a mixture of projectors $|0\rangle\langle 0|$ and $|1\rangle\langle 1|$, but the state is itself built by superposing $|\pm\rangle \equiv \frac{1}{\sqrt{2}}(|0\rangle \pm |1\rangle)$ and $|\pm i\rangle \equiv \frac{1}{\sqrt{2}}(|0\rangle \pm i|1\rangle)$, the two mutually unbiased bases w.r.t. the computational one. Therefore, her postselection cannot be used to distinguish which state has been prepared from the maximally mixed state $\mathbb{1}/2$. That is, she cannot deterministically distinguish states ensembles by choosing a suitable $p \neq 0$, which explains how the theory is operationally equivalent (sometimes called ‘tomographically indistinguishable’) to a qubit quantum theory with one ‘forbidden’ axis of the Bloch sphere. In other words, it is also guaranteed to be a perfectly well-behaved theory, in the sense that no matter the choice of V , N and its resolution $\{N_i\}$, no incoherence like overnormalized or negative probabilities can be found in the outcomes.

This one-party example is of course trivial, but as soon as more than one party is allowed, the possibility of deterministically sending a signal from one party to another will coincide with this kind of allowed postselection. This fact can be compared to the motivating example, where the choice of *which* quantum channel the elements of an instrument sum up to also amounts to deterministically inducing a bias in the outcome probability.

F.1 Bipartite biased quantum states

In the biased case, where $\overline{\mathcal{A}} = \overline{\mathcal{B}} = \text{Span}\{\mathbb{1}, \sigma_z\}$ the difference between a bipartite no signaling state and a general bipartite one becomes striking. The set of all valid states normalized on the local effects resolving $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$ is $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$ which, according to Proposition 1 is made of the following

13 basis elements:

$$\overline{\mathcal{A}} \otimes \overline{\mathcal{B}} \subset \text{Span}\left\{\mathbb{1}^A \otimes \mathbb{1}^B, \mathbb{1}^A \otimes \sigma_x^B, \mathbb{1}^A \otimes \sigma_y^B, \sigma_x^A \otimes \mathbb{1}^B, \sigma_x^A \otimes \sigma_x^B, \sigma_x^A \otimes \sigma_y^B, \sigma_x^A \otimes \sigma_z^B, \sigma_y^A \otimes \mathbb{1}^B, \sigma_y^A \otimes \sigma_x^B, \sigma_y^A \otimes \sigma_y^B, \sigma_y^A \otimes \sigma_z^B, \sigma_z^A \otimes \sigma_x^B, \sigma_z^A \otimes \sigma_y^B\right\}. \quad (266)$$

On the other hand, the no signaling composition of $\mathcal{A}$ and $\mathcal{B}$, $\mathcal{A} \otimes \mathcal{B}$ as by Definition 5, is only made of 9 elements (notice that the spanning set of $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$ is indeed contained in the one of $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$, in accordance with Definition 7). The four missing elements are

$$\overline{\mathcal{A}} \otimes \overline{\mathcal{B}} \setminus \mathcal{A} \otimes \mathcal{B} \subset \text{Span}\left\{\sigma_x^A \otimes \sigma_z^B, \sigma_y^A \otimes \sigma_z^B, \sigma_z^A \otimes \sigma_x^B, \sigma_z^A \otimes \sigma_y^B\right\}, \quad (267)$$

which are exactly the elements containing a σ_z term that we will now show to allow signaling between A and B . That the elements of $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}} \setminus \mathcal{A} \otimes \mathcal{B}$ allow for signaling is the content of Sec. 5; this result is proven in the general case in Lemma 5. Indeed, these basis elements are those which are quasi-orthogonal globally, meaning that an operator W that contains some of them will satisfy $\forall N^A \in \overline{\mathcal{A}}, \forall N^B \in \overline{\mathcal{B}}$

$$\begin{aligned} \text{Tr}\left[\left(N^A \otimes N^B\right) \cdot W\right] \\ = \frac{1}{d_A d_B} \text{Tr}\left[\left(N^A \otimes N^B\right)\right] \text{Tr}[W], \end{aligned} \quad (268)$$

because it satisfies Proposition (1) and therefore Eq. (17). But it does not obey quasi-orthogonality with respect to a local measurement, meaning that it will fail to satisfy the analog of Eq. (48) in at least one of Eqs. (61) like e.g., $\text{Tr}_A\left[\left(N^A \otimes \mathbb{1}\right) \cdot W\right] \neq 1/d_A \text{Tr}_A\left[N^A\right] \text{Tr}_A[W]$.

The consequence of this observation is that if Alice and Bob share a bipartite no signaling state, they may observe nonlocal entanglement effects on their outcome distributions, but these correlations will obey no signaling constraints in both directions, Eqs. (190). If, however, they share a general bipartite state, they may use it to achieve deterministic signaling: for certain states, they will be able to signal perfectly in one direction, i.e. Alice can perfectly send a message to Bob, and vice-versa.

For instance, consider the task where Alice receives a classical bit x and she wants to communicate it to Bob, so that his outcome b has the same value, $b = x$. Without any resources, Bob can only guess and thus succeed with $p(b=x)=1/2$. Now if we allow them to measure a shared state in $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$, they can pick the following state:

$$W_{A \prec B} = \frac{1}{4} \left(\mathbb{1}^A \otimes \mathbb{1}^B + \sigma_z^A \otimes \sigma_x^B \right), \quad (269)$$

(the notation in the subscript means ‘Alice can signal to Bob’) and choose to do the following: Alice ‘steers’ her measurement towards $|0\rangle$ or $|1\rangle$ depending on x ,

$$N_{|x}^A = \mathbb{1}^A + (-1)^x \sigma_z^A, \quad (270)$$

while Bob is measuring an unbiased $N^B = \mathbb{1}$ resolved into a measurement in the $|\pm\rangle$ basis,

$$N_b^B = \frac{1}{2} \left(\mathbb{1}^B + (-1)^b \sigma_x^B \right), \quad (271)$$

where $b = 0, 1$ so that his probabilistic effects satisfy $N_0^B + N_1^B = N^B = \mathbb{1}^B$. One can check that they are effectively properly normalized positive operators belonging to the proper state structures, $N_{|x}^A \in \mathcal{A}$, $N^B \in \mathcal{B}$, despite $N_0^B, N_1^B \notin \mathcal{B}$. The measurement yields the following probability distribution:

$$p(b|x) = \text{Tr} \left[\left(N_{|x}^A \otimes N_b^B \right) \cdot W_{A \prec B} \right]. \quad (272)$$

Injecting the above expressions into it yields

$$p(b|x) = \frac{1}{2} \left(1 + (-1)^{x+b} \right), \quad (273)$$

which gives 0 when $x \neq b$ and 1 when $x = b$; Alice’s setting is perfectly correlated to Bob’s outcome. We conclude that a bit was perfectly sent from A to B , or that using the $W_{A \prec B}$ state as their resource, they obtained $p(x = b) = 1$.

The bottom line of this example is that there exist states in $\overline{\mathcal{A}} \otimes \overline{\mathcal{B}}$ which allow beating one of the no signaling constraints, Eqs. (190), with a probability of 1. In this regard, and with respect to every theory that can be characterized by the projective methods of this article, quantum theory plays a special role as it is the only one whose bipartite states normalized on no signaling effects are automatically no signaling. This fact is the deeper meaning of the inclusion $\overline{\mathcal{P}_A} \otimes \overline{\mathcal{P}_B} \supseteq \mathcal{P}_A \otimes \mathcal{P}_B$ of Eq. (33) becoming an

equality in the case $\mathcal{P}_A = \mathcal{P}_B = \mathcal{I}$ (Eq. (139)) which follows from Lemma 6).

Comparing the biased to the regular quantum theory, the bipartite quantum states do saturate the inclusion (see example E.2), which is why they are no signaling compared to bipartite biased states. Contrastingly, bipartite quantum channels (see example E.4), when seen as the composition of two channels, do not saturate it (E.4). This is why information can be passed from one side of bipartite channels to the other, and why the no signaling channels are a proper subset of them (E.4).

F.2 Biased quantum channels

The dynamics of the biased quantum theory can be postulated as the maps from the biased theory to itself. Proposition 2 characterizes these as the set $\mathcal{A} \rightarrow \mathcal{B}$ where the state structure of the input and output are the same, $\mathcal{A} = \mathcal{B} = \text{Span}\{\mathbb{1}, \sigma_x, \sigma_y\}$:

$$\begin{aligned} \mathcal{A} \rightarrow \mathcal{B} \subset \text{Span} \{ & \mathbb{1}^A \otimes \mathbb{1}^B, \mathbb{1}^A \otimes \sigma_x^B, \mathbb{1}^A \otimes \sigma_y^B, \\ & \sigma_x^A \otimes \sigma_x^B, \sigma_x^A \otimes \sigma_y^B, \sigma_y^A \otimes \sigma_x^B, \sigma_y^A \otimes \sigma_y^B, \\ & \sigma_z^A \otimes \mathbb{1}^B, \sigma_z^A \otimes \sigma_x^B, \sigma_z^A \otimes \sigma_y^B, \sigma_z^A \otimes \sigma_z^B \}. \end{aligned} \quad (274)$$

Here, the fact proven in Sec. 5.1 that a transformation between state structures is a composition that does not forbid signaling in any direction can be proven by an example. Suppose Alice and Bob are sharing the channel $M_{A \prec B} = \frac{1}{2} \left(\mathbb{1}^A \otimes \mathbb{1}^B + \sigma_x^A \otimes \sigma_x^B \right)$. Alice can perfectly signal to Bob by encoding her setting x in the σ_x basis, $V_{|x} = 1/2(\mathbb{1} + (-1)^x \sigma_x)$, and if Bob measures in the same basis, they effectively have a perfect single bit channel, $p(b = x) = 1$.

On the other hand, suppose they share the channel $M_{A \succ B} = \frac{1}{2} \left(\mathbb{1}^A \otimes \mathbb{1}^B + \sigma_z^A \otimes \sigma_z^B \right)$. Now it is Bob who can perfectly signal to Alice: Alice has to use an ancilla so that she can prepare the same joint state as the bipartite example, Eq. (269),

$$W_{A \prec A'} = \frac{1}{4} \left(\mathbb{1}^A \otimes \mathbb{1}^{A'} + \sigma_z^A \otimes \sigma_x^{A'} \right). \quad (275)$$

She sends the A part through the channel and keeps the A' part in her ancilla. Bob can then apply the measurement $N_{|y}^B = \mathbb{1} + (-1)^y \sigma_z$ at the outcome of the channel, depending on the

variable y he wishes to send. Alice can then finally measure her ancilla in the σ_x basis, $N_a^{A'} = 1/2(\mathbb{1} + (-1)^a \sigma_x)$, leading to the distribution

$$\begin{aligned} p(a|y) \\ = \text{Tr} \left[\left(N_{|y}^B \otimes N_a^{A'} \right) \cdot \text{Tr}_A [M_{A \succ B} \cdot W_{A \prec A'}]^T \right] . \end{aligned} \quad (276)$$

and she will get perfect correlation with Bob setting, $p(a = y) = 1$, exactly like in the bipartite state example. Therefore, in the state structure $\mathcal{A} \rightarrow \mathcal{B}$, there are elements allowing perfect signaling from Alice to Bob as well as from Bob to Alice. Contrasingly, the quantum channels (seen as a composition of an input and an output using the transformation) only allow one direction of signaling, from the input (Alice) to the output (Bob; see Remark E.3).

G Introductory example extended: dynamics-like construction up to fourth order

For this example, we will focus on the case of higher-order transformations built upon quantum mechanics. In a nutshell, we will go back to the construction that lifts the POVM formalism into the quantum instrument formalism that we presented at the beginning of Sec. 3, we will abstract it, and we will repeat it until an indefinite causal order (ICO) arises. The reason we do that is two-fold: on the one hand, we want to present how the characterization techniques work in a concrete case (as we will see, all the objects that will be defined have already been studied in the literature). On the other hand, and in accordance with the discussion of Sec. 6.2, we want to stress how tame quantum theory is with respect to other theories that can be characterized using projective means: albeit we will be using the transformation operation $\rightarrow$ repetitively, i.e. we are nesting operations in a way that allows for bidirectional signaling, it will not result in a transformation with ICO before the fourth order. This is in stark contrast with the example of the biased theory of App. F in which one could observe signaling in both directions already at the level of bipartite states as well as of channels.

G.1 Reformulating the introductory example: the quantum channel is a 1-comb

Reformulating the introductory example of Sec. 3, we start with a state structure $\mathcal{A}_0 \subset \mathcal{L}(\mathcal{H}^{A_0})$ with projector $\mathcal{I}_{A_0}$ and a trace of 1, as in Eqs. (8) or (234). Defining the functionals on this state structure, we consider state and effect pairs in complementary (i.e. quasi-orthogonal) state structures, $(\rho, \mathbb{1}) \in \mathcal{A}_0 \times \overline{\mathcal{A}_0}$, linked by the normalization of the probability rule

$$1 = \langle \mathbb{1}, \rho \rangle_{\mathcal{L}(\mathcal{H}^{A_0})} \equiv \text{Tr} [\mathbb{1} \cdot \rho] . \quad (277)$$

These pairs represent all deterministic preparation and measurement procedures, i.e. those yielding a probability of 1 irrespective of the choice of state. These pairs characterize the ‘first-order theory’ (because these are the states and effects on which can be defined first-order transformations, which in turn will be the effects of the second-order theory). Probabilistic assignments are obtained by resolving the effect state structure by a collection of positive operators $\{E_i\}$,

$$p(i|\rho, \mathbb{1}) = \text{Tr} [E_i \cdot \rho] , \quad (278)$$

yielding the POVM formalism and the usual definition of effects.

To go to the second-order theory, one postulates some dynamics so that the state structure $\mathcal{A}_0$ is mapped to a similarly defined state structure $\mathcal{A}_1$ by first-order transformations, i.e. by some CPTP map $\mathcal{M} \in \mathcal{L}(\mathcal{L}(\mathcal{H}^{A_0}), \mathcal{L}(\mathcal{H}^{A_1}))$ so that

$$\langle \mathbb{1}, \rho \rangle_{\mathcal{L}(\mathcal{H}^{A_0})} \mapsto \langle \mathbb{1}, \mathcal{M}^{A_0 \rightarrow A_1}(\rho) \rangle_{\mathcal{L}(\mathcal{H}^{A_1})} . \quad (279)$$

As it was shown in example E.3, in the CJ picture, this results in a new pair $(\rho^{A_0} \otimes \mathbb{1}^{A_1}, M^{A_0 A_1}) \in (\mathcal{A}_0 \otimes \overline{\mathcal{A}_1}) \times (\mathcal{A}_0 \rightarrow \mathcal{A}_1)$, so that

$$1 = \text{Tr} \left[M \cdot (\rho \otimes \mathbb{1}^T) \right]^T = \text{Tr} [M \cdot (\rho \otimes \mathbb{1})] . \quad (280)$$

Here, M is a quantum channel in CJ representation, which by definition is a 1-comb. Resolving state structure $(\mathcal{A}_0 \rightarrow \mathcal{A}_1)$ yields the quantum instrument formalism,

$$p(j|M, \rho) = \langle M_j, \rho \otimes \mathbb{1} \rangle . \quad (281)$$

In this example, the dynamics-like construction of a second-order theory is completed by requiring

that the set of states be extended to all functionals normalized on the effects. That is, to all W in $\overline{\mathcal{A}}_0 \rightarrow \overline{\mathcal{A}}_1$ so that the normalization becomes

$$\langle M, W \rangle_{\mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})} = 1, \quad (282)$$

where W is a single-partite PM and M is a 1-comb, so that both are second-order objects.

Nonetheless, we noted at the end of example E.3 that single-partite PM trivially decomposes into states and measurements. Meaning that the states W of the second-order naturally descend to the first-order: $\overline{\mathcal{A}}_0 \rightarrow \overline{\mathcal{A}}_1 = \mathcal{A}_0 \otimes \overline{\mathcal{A}}_1$. The explanation in terms of the projector algebra is simply a matter of definition $\mathcal{P}_{A_0} \rightarrow \mathcal{P}_{A_1} = \overline{\mathcal{P}}_{A_0} \otimes \overline{\mathcal{P}}_{A_1}$. In terms of signaling structure, because the transformation is two-way signaling, $\mathcal{P}_{A_0} \rightarrow \mathcal{P}_{A_1} = (\overline{\mathcal{P}}_{A_0} \prec \mathcal{P}_{A_1}) \cup (\mathcal{P}_{A_0} \succ \overline{\mathcal{P}}_{A_1})$, the functional on transformations have to be no signaling by De Morgan duality (42c): $\overline{\mathcal{P}}_{A_0} \rightarrow \overline{\mathcal{P}}_{A_1} = \overline{\mathcal{P}}_{A_0} \prec \mathcal{P}_{A_1} \cap \overline{\mathcal{P}}_{A_0} \succ \mathcal{P}_{A_1} = (\overline{\mathcal{P}}_{A_0} \prec \overline{\mathcal{P}}_{A_1}) \cap (\mathcal{P}_{A_0} \succ \overline{\mathcal{P}}_{A_1}) \stackrel{(56a)}{=} \mathcal{P}_{A_0} \otimes \overline{\mathcal{P}}_{A_1}$.

Another difference we noticed in Remark E.3 is that the 1-comb was one-way signaling despite being built by transformation. Now we can see this fact quickly from isomorphism (183) (Eq. (76a) in the main text): observe that the 1-comb is characterized by $\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1}$ which by Lemma 6 is equivalent to

$$\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_1} = \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1}. \quad (283)$$

Hence, we are guaranteed that the 1-combs have a fixed signaling direction, so they cannot show indefinite causal order. If the base state structure was anything else than quantum states, we would not have witnessed such a simplification; the 1-comb projector would have decomposed into two different orderings: $\mathcal{P}_{A_0} \rightarrow \mathcal{P}_{A_1} = (\overline{\mathcal{P}}_{A_0} \prec \mathcal{P}_{A_1}) \cup (\overline{\mathcal{P}}_{A_0} \succ \mathcal{P}_{A_1})$.

G.2 Dynamics-like construction

As preparation and measurement are special transformations with respectively, input and output trivial systems, we can also promote the destructive measurement at A_1 into a measurement performed by a second party, so that the probability rule becomes

$$p(i, j | M, \rho) = \langle M_j, \rho \otimes E_i^T \rangle, \quad (284)$$

this is still one-way signaling as it amounts to having a second party making a POVM measurement $\{E_i\}$ (which is a special case of a quantum instrument) after the quantum instrument $\{M_j\}$.

Remark. To better stick to our conventions, we could have passed all probabilistic assignments to the left side of the inner product:

$$p(i, j | M, \rho) = \langle E_i * M_j, \rho \rangle_{\mathcal{L}(\mathcal{H}^{A_0})}, \quad (285)$$

where $E_i * M_j \equiv \text{Tr}_{A_1} \left[\left(\mathbb{1}^{A_0} \otimes E_i^T \right) \cdot M_j \right]$ is the link product of the two operators [3], which corresponds to the CJ representation of the sequential composition of these two instruments over subsystem A_1 .

This construction of the quantum instrument formalism has been abstracted under the name ‘dynamics-like construction’ at the end of App. E.5. For the sake of the argument, however, we will take a slightly different point of view in the following: the introduced transformations will be considered to be the states of the higher-order theory instead of the effects, and the effects will not be extended to be the functionals on the new states. In other words, the introduced dynamics will be considered to be the deterministic objects over which the parties have no influence (the state), and everything else will be locally controlled by some parties acting probabilistically (the effects). For the case of Eq. (284), this amounts to taking the marginal over j , $p(i | M, \rho) = \langle M, \rho \otimes E_i^T \rangle$. The parties A_0 and A_1 then share a channel in which A_0 inputs a state, and of which A_1 measures the output. We do so because this will maximize the number of local parties, which makes ICO more obvious: it would be difficult to talk about the ICO *within* the whole operator representing the shared dynamics (i.e. the global functional on the state), whereas it can be more easily discussed by looking at the correlations achievable by a group of parties sharing this global object. In this regard, this amounts to taking the dynamics as the state and everything else as local effects in tensor product.

Therefore, given states in $\mathcal{A} \in \mathcal{L}(\mathcal{H}^{A_0})$ and effects in $\overline{\mathcal{A}}$, the dynamics-like construction now amounts to defining bipartite states in $\mathcal{A} \rightarrow \mathcal{B} \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_1})$ with $\mathcal{H}^{A_1} \cong \mathcal{H}^{A_0}$. Associated with these shared states are two local effects $\mathcal{A}_0$

and $\overline{\mathcal{A}}_1$, each under the control of an independent party. I.e., the dynamics-like construction of the next order consists of the following redefinition of the states and effects pair:

$$\mathcal{A}_0 \times \overline{\mathcal{A}}_0 \mapsto (\mathcal{A}_0 \rightarrow \mathcal{A}_1) \times (\overline{\mathcal{A}}_0 \rightarrow \overline{\mathcal{A}}_1) . \quad (286)$$

G.3 The quantum supermap is a 2-comb

The third-order theory is obtained by assuming the existence of dynamics over the current dynamics, i.e. second-order transformations. In the same way that structure-preserving maps can be nicknamed ‘superoperator’, we are introducing a ‘supermap’ $\mathcal{N}$ as a linear map between two maps with the same state structure [2]. If the map $\mathcal{M}$ of the previous section is defined between subsystems A_1 and A_2 , we introduce $\mathcal{N}$ as a CPTP-preserving supermap that send $\mathcal{M}$ to a similar map $\tilde{\mathcal{M}}$ between subsystems A_0 and A_3 , $\mathcal{N}(\mathcal{M}^{A_1 \rightarrow A_2}) = \tilde{\mathcal{M}}^{A_0 \rightarrow A_3}$, so that

$$\langle \mathbb{1}, \mathcal{M}(\rho) \rangle_{\mathcal{L}(\mathcal{H}^{A_1})} \mapsto \langle \mathbb{1}, [\mathcal{N}(\mathcal{M})](\rho) \rangle_{\mathcal{L}(\mathcal{H}^{A_3})} . \quad (287)$$

Going to the CJ picture, the following probability rule is obtained:

$$p(i, j | N, M, \rho) = \langle N, \rho \otimes M_j^{T_{A_2}} \otimes \overline{E}_i^T \rangle . \quad (288)$$

The theory is then characterized by quasi-orthogonal pairs of the form $(N, \rho \otimes M \otimes \mathbb{1}) \in ((\mathcal{A}_1 \rightarrow \mathcal{A}_2) \rightarrow (\mathcal{A}_0 \rightarrow \mathcal{A}_3)) \times (\mathcal{A}_0 \otimes (\mathcal{A}_1 \rightarrow \mathcal{A}_2) \otimes \overline{\mathcal{A}}_3)$.

This way, N is a 2-comb, an object that transforms 1-combs into 1-combs. We have already seen in the last section that the 1-combs themselves can be interpreted as one-way signaling objects, so that the signaling of the 2-comb can be made explicit by looking at its corresponding projector:

$$\begin{aligned} & (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_2}) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_3}) \\ & \stackrel{(183)}{=} (\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2}) \rightarrow (\overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3}) \\ & \stackrel{(166)}{=} (\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2}) \prec (\overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3}) \\ & \cup (\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2}) \succ (\overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3}) \\ & = \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3} \\ & \cup \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} . \quad (289) \end{aligned}$$

So it may be concluded that the most general 2-comb is a superposition of two possible quantum networks: one in which the channel M is

measured, $\mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2}$, then reprepared as $\tilde{M}$, $\overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3}$ in its causal future and one where it is first $\tilde{M}$ which is prepared.

Yet, there is again an isomorphism at play. This is the isomorphism presented in Sec. 6.2: the n -combs are equivalent to quantum networks. This will make the interpretation of the supermap N , a second-order transformation, equivalent to a succession of first-order transformations with a single signaling direction.

$$\begin{aligned} & (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_2}) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_3}) \\ & \stackrel{(144)}{=} (\mathcal{I}_{A_0} \otimes (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_2})) \rightarrow \mathcal{I}_{A_3} \\ & \stackrel{(183)}{=} (\overline{\mathcal{I}}_{A_0} \otimes (\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2})) \prec \mathcal{I}_{A_3} \\ & \stackrel{(186)}{=} (\overline{\mathcal{I}}_{A_0} \prec (\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2})) \prec \mathcal{I}_{A_3} \\ & = \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_3} . \quad (290) \end{aligned}$$

This means that N can equivalently be understood as a succession of two quantum channels [3, 4, 5]; this is a peculiar feature of quantum supermaps, and of quantum combs in general.

This surprisingly tame behavior of having a single signaling direction when reduced to a succession of first-order objects will not be carried out in the fourth order, as we will see next. In addition, note that this explains why a single-partite PM does not feature ICO (or anything outside of the usual quantum theory). If the systems A_0 and A_3 were trivial, so that N is defined as (the CJ of) a supermap that takes a quantum channel to a probability, we see that the associated projector $\overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_3} \mapsto 1 \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec 1$, and this can be simplified into $\mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \stackrel{(186)}{=} \mathcal{I}_{A_1} \otimes \mathcal{D}_{A_2}$.

G.4 The quantum super-supermap is an MPM

To go to the fourth order, the same procedure is applied again: we introduce a ‘super-supermap’ $\mathcal{W}$ so that it maps supermaps like $\mathcal{N}$ to themselves, $\mathcal{W}(\mathcal{N}) = \tilde{\mathcal{N}}$, where $\mathcal{N} \in \mathcal{L}(\mathcal{H}^{A_1} \otimes \mathcal{H}^{A_2} \otimes \mathcal{H}^{A_5} \otimes \mathcal{H}^{A_6})$ and $\tilde{\mathcal{N}} \in \mathcal{L}(\mathcal{H}^{A_0} \otimes \mathcal{H}^{A_3} \otimes \mathcal{H}^{A_4} \otimes \mathcal{H}^{A_7})$:

$$\begin{aligned} & \langle \mathbb{1}, [\mathcal{N}(\mathcal{M})](\rho) \rangle_{\mathcal{L}(\mathcal{H}^{A_3})} \\ & \mapsto \langle \mathbb{1}, [[\mathcal{W}(\mathcal{N})](\mathcal{M})](\rho) \rangle_{\mathcal{L}(\mathcal{H}^{A_7})} . \quad (291) \end{aligned}$$

In the CJ representation, the probability rule is

$$p(i, j, k | W, M, N, \rho) = \langle W, \rho \otimes N_j^{T_{A_2} T_{A_6}} \otimes M_j^{T_{A_4}} \otimes \bar{E}_i^{T_{A_7}} \rangle. \quad (292)$$

Here, M is a transformation of ρ defined in $\mathcal{A}_3 \rightarrow \mathcal{A}_4$; N is a transformation of M , $N \in (\mathcal{A}_2 \rightarrow \mathcal{A}_5) \rightarrow (\mathcal{A}_1 \rightarrow \mathcal{A}_6)$, and the W operator is a transformation of N ,

$$W \in [(\mathcal{A}_2 \rightarrow \mathcal{A}_5) \rightarrow (\mathcal{A}_1 \rightarrow \mathcal{A}_6)] \rightarrow [(\mathcal{A}_3 \rightarrow \mathcal{A}_4) \rightarrow (\mathcal{A}_0 \rightarrow \mathcal{A}_7)]. \quad (293)$$

We can now treat N as a party itself, which physically corresponds to someone acting a first time between nodes A_1 and A_2 , and a second time between A_5 and A_6 , using a quantum network.

Again, one can focus on the projector characterizing the state structure of N to extract its signaling structure by putting it in normal form. First, successive applications of the uncurrying rule (144) yield

$$\begin{aligned} & [(\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6})] \\ & \rightarrow [(\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_7})] \\ & = [(\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \otimes ((\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \\ & \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6}))] \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_7}) \\ & = [\mathcal{I}_{A_0} \otimes ((\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \\ & \otimes ((\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6})))] \rightarrow \mathcal{I}_{A_7}. \end{aligned} \quad (294)$$

Next, Eqs. (183) and (290), (186), (183), and (161) are used successively:

$$\begin{aligned} & [\mathcal{I}_{A_0} \otimes ((\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \\ & \otimes ((\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6})))] \rightarrow \mathcal{I}_{A_7} \\ & = [\mathcal{I}_{A_0} \otimes ((\bar{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4}) \\ & \otimes (\bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \bar{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6}))] \rightarrow \mathcal{I}_{A_7} \\ & \stackrel{(186)}{=} [\mathcal{I}_{A_0} \prec ((\bar{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4}) \\ & \otimes (\bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \bar{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6}))] \rightarrow \mathcal{I}_{A_7} \stackrel{(183)}{=} \\ & \overline{\mathcal{I}_{A_0} \prec ((\bar{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4}) \otimes (\bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \bar{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6}))} \\ & \prec \mathcal{I}_{A_7} \stackrel{(161)}{=} \\ & \overline{\bar{\mathcal{I}}_{A_0} \prec ((\bar{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4}) \otimes (\bar{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \bar{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6}))} \\ & \prec \mathcal{I}_{A_7}. \end{aligned} \quad (295)$$

If, for simplicity, we assume the subsystems A_0 and A_7 to be of dimension 1, we can already notice in the first line of the previous equation that the state structure of the operator W features a projector whose expression involves the negation of a tensor product of two quantum combs:

$$W \in \overline{(\mathcal{A}_2 \rightarrow \mathcal{A}_5) \rightarrow (\mathcal{A}_1 \rightarrow \mathcal{A}_6) \otimes (\mathcal{A}_3 \rightarrow \mathcal{A}_4)}, \quad (296)$$

i.e. it is the functional normalized on the tensor product between a 2-comb and 1-comb. This is by definition a multi-round process matrix (MPM). We proved in a previous work that it can manifest indefinite causal order and even more that it can beat causal inequalities [11]. The fourth order, therefore, presents multiple directions of signaling in a non-trivial manner.

The fundamental reason why, as we showed in this previous article [11], is that this last expression defines a set which is the union of 3 sets of quantum combs with different signaling orderings,

$$\begin{aligned} & \overline{(\mathcal{A}_2 \rightarrow \mathcal{A}_5) \rightarrow (\mathcal{A}_1 \rightarrow \mathcal{A}_6) \otimes (\mathcal{A}_3 \rightarrow \mathcal{A}_4)} \\ & = \\ & \overline{(((\mathcal{A}_1 \rightarrow \mathcal{A}_2) \rightarrow \mathcal{A}_5) \rightarrow \mathcal{A}_6) \rightarrow \mathcal{A}_3) \rightarrow \mathcal{A}_4 \cup} \\ & \overline{(((\mathcal{A}_1 \rightarrow \mathcal{A}_2) \rightarrow \mathcal{A}_3) \rightarrow \mathcal{A}_4) \rightarrow \mathcal{A}_5) \rightarrow \mathcal{A}_6 \cup} \\ & \overline{(((\mathcal{A}_3 \rightarrow \mathcal{A}_4) \rightarrow \mathcal{A}_1) \rightarrow \mathcal{A}_2) \rightarrow \mathcal{A}_5) \rightarrow \mathcal{A}_6}. \end{aligned} \quad (297)$$

In other words, W can be taken as a superposition of all combs that respect the signaling ordering in between the nodes of the 2-comb and the 1-comb that will be plugged into it, but it itself does not assume a global ordering. The above expression thus has 3 possible global orderings, depending on whether the operation of the party with the 1-comb (acting between nodes 3 and 4) is before, in between, or after the two operations of the party with the 2-comb.

Back to the general case of a super-supermap, we can use this insight to present a case of non-unique normal form. One can further simplify the

projector to a normal form using Eq. (56a):

$$\begin{aligned}
& \overline{\mathcal{I}}_{A_0} \prec \\
& \overline{\left(\left(\overline{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4} \right) \otimes \left(\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \overline{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6} \right) \right)} \\
& \quad \prec \mathcal{I}_{A_7} \\
& \quad = \overline{\mathcal{I}}_{A_0} \prec \\
& \overline{\left(\left(\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \overline{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6} \prec \overline{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4} \right) \cap \right.} \\
& \quad \left. \left(\overline{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4} \prec \overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \overline{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6} \right) \right) \\
& \quad \prec \mathcal{I}_{A_7} \\
& \quad \stackrel{(119b)}{=} \overline{\mathcal{I}}_{A_0} \prec \\
& \overline{\left(\left(\overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \overline{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6} \prec \overline{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4} \right) \cup \right.} \\
& \quad \left. \left(\overline{\mathcal{I}}_{A_3} \prec \mathcal{I}_{A_4} \prec \overline{\mathcal{I}}_{A_1} \prec \mathcal{I}_{A_2} \prec \overline{\mathcal{I}}_{A_5} \prec \mathcal{I}_{A_6} \right) \right) \prec \mathcal{I}_{A_7} \\
& \quad \stackrel{(66)}{=} \overline{\mathcal{I}}_{A_0} \prec \\
& \quad \left(\left(\mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \right) \cup \right. \\
& \quad \left. \left(\mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \right) \right) \prec \mathcal{I}_{A_7} \\
& \quad \stackrel{(69)}{=} \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \\
& \quad \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_7} \cup \\
& \quad \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \\
& \quad \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_7} . \quad (298)
\end{aligned}$$

Thus,

$$\begin{aligned}
& [(\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6})] \\
& \rightarrow [(\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_7})] \\
& \quad = \\
& \quad \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_1} \\
& \quad \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_7} \\
& \quad \cup \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \\
& \quad \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_7} . \quad (299)
\end{aligned}$$

But one can also inject the previously discussed simplification for the case of a functional into the central part. Distributing over the union, one

obtains a normal form as in Definition 9:

$$\begin{aligned}
& [(\mathcal{I}_{A_2} \rightarrow \mathcal{I}_{A_5}) \rightarrow (\mathcal{I}_{A_1} \rightarrow \mathcal{I}_{A_6})] \\
& \rightarrow [(\mathcal{I}_{A_3} \rightarrow \mathcal{I}_{A_4}) \rightarrow (\mathcal{I}_{A_0} \rightarrow \mathcal{I}_{A_7})] \\
& \quad = \\
& \quad \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_5} \\
& \quad \quad \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_7} \\
& \quad \cup \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \\
& \quad \quad \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_7} \\
& \quad \cup \overline{\mathcal{I}}_{A_0} \prec \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2} \prec \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6} \prec \mathcal{I}_{A_3} \\
& \quad \quad \prec \overline{\mathcal{I}}_{A_4} \prec \mathcal{I}_{A_7} . \quad (300)
\end{aligned}$$

The reason for these two normal forms to be valid comes from the transitivity of the prec, i.e. that it is associative, Eq. (67). Call $\mathcal{P}_A = \mathcal{I}_{A_3} \prec \overline{\mathcal{I}}_{A_4}$, $\mathcal{P}_{B_0} = \mathcal{I}_{A_1} \prec \overline{\mathcal{I}}_{A_2}$, and $\mathcal{P}_{B_1} = \mathcal{I}_{A_5} \prec \overline{\mathcal{I}}_{A_6}$. Then, in the expression $(\mathcal{P}_A \prec \mathcal{P}_{B_0} \prec \mathcal{P}_{B_1}) \cup (\mathcal{P}_{B_0} \prec \mathcal{P}_{B_1} \prec \mathcal{P}_A)$ one sees that the state structure characterized by this projector allows for signaling from A to B_0 and B_1 and at the same time from B_0 and B_1 to A . Obviously, if A can be both before or after the two B 's, she can also be in the middle, in equation:

$$\begin{aligned}
& (\mathcal{P}_A \prec \mathcal{P}_{B_0} \prec \mathcal{P}_{B_1}) \cup (\mathcal{P}_{B_0} \prec \mathcal{P}_{B_1} \prec \mathcal{P}_A) = \\
& (\mathcal{P}_A \prec \mathcal{P}_{B_0} \prec \mathcal{P}_{B_1}) \cup (\mathcal{P}_{B_0} \prec \mathcal{P}_{B_1} \prec \mathcal{P}_A) \\
& \quad \cup (\mathcal{P}_{B_0} \prec \mathcal{P}_A \prec \mathcal{P}_{B_1}) . \quad (301)
\end{aligned}$$

In other words, saying that B_0 can signal to A which can signal to B_1 does not add new information when one knows that A can signal to B_0 and B_1 and at the same time that B_0 and B_1 can signal to A . This implies by De Morgan duality that

$$\begin{aligned}
& (\mathcal{P}_A \prec \mathcal{P}_{B_0} \prec \mathcal{P}_{B_1}) \cap (\mathcal{P}_{B_0} \prec \mathcal{P}_{B_1} \prec \mathcal{P}_A) = \\
& (\mathcal{P}_A \prec \mathcal{P}_{B_0} \prec \mathcal{P}_{B_1}) \cap (\mathcal{P}_{B_0} \prec \mathcal{P}_{B_1} \prec \mathcal{P}_A) \\
& \quad \cap (\mathcal{P}_{B_0} \prec \mathcal{P}_A \prec \mathcal{P}_{B_1}) , \quad (302)
\end{aligned}$$

and thus that Eqs. (299) and (300) are equivalent. Ergo, the normal form is multiply defined in that case because of redundant information in the description of the signaling structure.

Redundancy or not, the presence of a union in the normal form indeed confirms that more than one fixed signaling direction is allowed, hence that there are operators in that state structure that allow for indefinite causal ordering. What we have just done is to track the ICO origin as an ambiguity in the decomposition of higher-order into

lower-order: using the algebraic rules of the projectors, we could express the projector associated with the fourth-order in terms of unions of first-order projectors (by that we mean that we expressed the state structure of super-supermap as unions of ‘prec chains’ between quantum states and measurements). This decomposition does not forbid the existence of operators that belong to the affine hull of the union. That is, operators that are obtained from a superposition of several chains with different signaling directions. These are the ones that can show causal non-separability, like the operator studied in Ref. [11] for example.

The bottom line of this example is that we had to go to the fourth order to observe signaling in more than one direction and an indefinite causal order. In a general process theory built using the dynamics-like construction, there should have been two directions already at the second order, possibly featuring ICO (see the biased quantum theory example), and 2^{n-1} directions at order n . This tameness is a concrete consequence of Lemma 6 and the ensuing Theorem 3.