

Directed Univalence: From Riehl–Shulman to a Complete Principle

YonedaAI Research
Magneton Foundational Studies

4 May 2026

Abstract

Voevodsky’s univalence axiom equates path identification with type-equivalence and lies at the heart of homotopy type theory (HoTT). Its directed analogue — which would equate hom-types in the universe with functions between types — has been pursued since the introduction of *simplicial type theory* (STT) by Riehl and Shulman (2017). Riehl–Shulman gives a synthetic theory of $(\infty, 1)$ -categories whose internal hom-types $\mathrm{Hom}_A(a, b)$ behave correctly inside Segal types, and a *partial* directed univalence holds for covariant fibrations: the universe of discrete types restricted along a covariant family is recovered from the family itself. A complete principle for the entire universe remained out of reach. Gratzer, Weinberger and Buchholtz (2024) introduce *triangulated type theory* (TTT), an extension of STT with a system of modalities ($\flat$, $\sharp$, Disc , $\mathrm{opposite}$), and prove a directed univalence theorem for the universe $\mathcal{S}$ of *discrete* types: $\mathrm{Hom}_{\mathcal{S}}(A, B) \simeq \mathrm{Fun}(A, B)$. This is the first complete fragment of a directed univalence principle.

We give an expository, structured account of the line of work from Riehl–Shulman 2017 through Gratzer–Weinberger–Buchholtz 2024. We isolate the technical contents of partial directed univalence, formulate four candidate *full* directed univalence principles of increasing strength, and analyse what would be required to prove each. We propose a programme of *layered* directed univalence indexed by the Segal/Rezk/discrete hierarchy, give toy implementations in Haskell and Lean 4, and connect the project to natural number objects in $(\infty, 1)$ -toposes, higher inductive types, and synthetic representation theory. The complete principle remains open; we make explicit the obstructions, including modal dependence, opposite-type coherence, and the lack of an internal universe of $(\infty, 1)$ -categories.

Contents

1	Introduction	3
2	Background: Symmetric Univalence and Directed Motivation	4
2.1	Voevodsky’s univalence axiom	4
2.2	What “directed” should mean	5
2.3	$(\infty, 1)$ -categorical motivation	6

3	Riehl–Shulman Simplicial Type Theory	6
3.1	Topes, cubes, extension types	6
3.2	Hom-types and Segal types	7
3.3	Discrete types	7
3.4	Synthetic Yoneda lemma	8
4	Partial Directed Univalence (Riehl–Shulman 2017)	8
4.1	Covariant fibrations	8
4.2	What is partial about it	8
4.3	Limitations made formal	9
5	Triangulated Type Theory	9
5.1	Modal layer	9
5.2	The discrete universe $\mathcal{S}$	10
5.3	The directed univalence theorem	10
5.4	Worked examples in the discrete universe	10
5.5	Comparison to the symmetric universe	11
5.6	Implementation status in <code>rzk</code>	11
5.7	What is and is not proved	12
6	Toward a Complete Principle	12
6.1	Four candidates	12
6.2	Layered programme	12
6.3	Obstructions to Candidates II–IV	13
6.4	Sketch of a proof attempt for Candidate III	13
6.5	The Cisinski–Nguyen line of work	14
6.6	Cubical comparison	14
6.7	An alternative axiomatic approach	14
6.8	Summary table	15
7	Connections to $(\infty, 1)$-Natural Number Objects and HITs	15
7.1	Natural number objects in directed type theory	15
7.2	Higher inductive types and directed univalence	15
7.3	Synthetic representation theory	16
7.4	Synthetic adjunctions and Kan extensions	16
7.5	Stability under base change	16
7.6	Comparison to formalisations: <code>Coq-HoTT</code> , <code>Lean-mathlib</code> , <code>Cubical Agda</code> , <code>rzk</code>	17
8	Discussion	17
8.1	Why is directed univalence so hard?	17
8.2	What would change if Candidate III were proven?	17
8.3	What would change if Candidate IV were proven?	17
8.4	Limitations of this paper	18
9	Open Problems	18

1 Introduction

Univalence is the structural principle that distinguishes homotopy type theory from intensional Martin–Löf type theory. Its statement,

$$\text{ua} : (A \simeq B) \xrightarrow{\simeq} (A =_{\mathcal{U}} B),$$

internalises Mac Lane’s slogan *isomorphic objects are equal* into the type theory itself [1, 2]. Once one accepts univalence, virtually every homotopy-theoretic and category-theoretic phenomenon admits a natural type-theoretic expression: equivalences become identities, transport along an equivalence is automatic, and the structure identity principle becomes a theorem rather than a methodological wish.

Yet univalence as stated lives in a fundamentally *symmetric* world. The identity type $A =_{\mathcal{U}} B$ is a groupoid: it has inverses. This reflects a basic fact about the type-theoretic universe: every type, in HoTT, is treated as an ∞ -groupoid. Categories proper — where morphisms can be non-invertible — do not admit a clean internal description in HoTT without artefacts.

Directed type theory aims to repair this asymmetry by introducing a directed analogue of the path type, often written $\text{Hom}_A(a, b)$, which need not be invertible. The objective is a synthetic theory of $(\infty, 1)$ -categories analogous to the way HoTT serves as a synthetic theory of ∞ -groupoids. Riehl and Shulman’s 2017 paper [3] introduces simplicial type theory (**STT**) as such a foundation. The fundamental new ingredient is a strict interval $\mathbf{2}$ with two distinct endpoints $0, 1 : \mathbf{2}$ and a sub-type-system of *topes* (propositions over the cube of intervals) that lets one write *extension types* of the form

$$\langle \prod_{t:\mathbf{2}} A(t) \mid \Phi \mapsto a \rangle,$$

the type of dependent functions on $\mathbf{2}$ that on the tope Φ agree definitionally with a prescribed partial function a . Hom-types are then defined as extension types out of $\mathbf{2}$.

Within **STT** there is no *a priori* guarantee that an arbitrary type A describes an $(\infty, 1)$ -category — only that it has a mapping space structure. Two predicates $\text{isSegal}(A)$ (composition of hom-types) and $\text{isRezk}(A)$ (equivalences are identities) cut out the types of interest, the Rezk types, which model complete Segal spaces and hence $(\infty, 1)$ -categories.

The natural question is then: *is there a directed analogue of univalence?* If $\mathcal{S}$ is some directed universe of types, is the hom-type $\text{Hom}_{\mathcal{S}}(A, B)$ equivalent to the type of functions $\text{Fun}(A, B) := A \rightarrow B$? Riehl–Shulman prove a partial result: a covariant family $P : A \rightarrow \mathcal{U}$ is classified by a functor $A \rightarrow \mathcal{S}_{\text{disc}}$ into the *discrete* sub-universe, and the data of the family is recovered up to equivalence from the functor. This is directed univalence *for covariant families* but not for the universe at large.

For seven years, the question of whether the directed universe itself satisfies a directed univalence principle was open. Gratzer, Weinberger and Buchholtz (2024) settle the question

for the discrete sub-universe: they introduce a system of modalities adequate to express both the directed universe and the symmetric universe of ∞ -groupoids inside it, and prove

$$\mathbf{dua}_{\mathcal{S}} : \mathbf{Hom}_{\mathcal{S}}(A, B) \simeq \mathbf{Fun}(A, B),$$

where $\mathcal{S}$ is the universe of discrete types [4]. They name the resulting system *triangulated type theory* (**TTT**), reflecting the appearance of a triangle of modalities $\flat \dashv \mathbf{Disc} \dashv \sharp$ analogous to Schreiber–Shulman cohesion.

This paper has four aims:

1. Give a self-contained, structured account of **STT** (Section 3) and partial directed univalence (Section 4).
2. Present the Gratzer–Weinberger–Buchholtz construction (Section 5) and isolate exactly what is and is not proved.
3. Formulate four candidate *complete* directed univalence principles (Section 6) and analyse the obstructions.
4. Connect directed univalence to higher inductive types, $(\infty, 1)$ -natural number objects, and synthetic representation theory (Section 7).

We accompany the exposition with formal sketches in Haskell (under `src/`) and Lean 4 (under `lean/`) that mechanise the relevant equivalences for the discrete fragment and lay out the directed universe API. The Haskell modules are not Lean proofs — they implement *toy* models that exhibit the algebraic content (covariant fibrations as $A \rightarrow \mathcal{S}$ functors, hom-types as Day convolutions on a **2**-action) so that the conceptual content can be machine-checked even if foundationally weakened. Lean 4 sketches use postulates where appropriate.

Open problems are collected in Section 9. We list four: extension to non-discrete universes, $(\infty, 1)$ -Cat as a primitive type, internal language theorems for cocartesian fibrations, and the connection to Langlands.

Conventions. We write $\mathcal{U}$ for an ambient (symmetric) universe; $\mathcal{S}$ for a directed sub-universe; $A : \mathcal{U}$ or $A : \mathcal{S}$ for a type/object; $a, b : A$ for terms. The directed interval is $\mathbf{2} = \{0 \leq 1\}$. Hom-types are written $\mathbf{Hom}_A(a, b)$ and identity types $a =_A b$. We use $f \simeq g$ for the type of equivalences. Throughout, **STT** refers to Riehl–Shulman 2017 simplicial type theory, **TTT** to Gratzer–Weinberger–Buchholtz 2024 triangulated type theory.

2 Background: Symmetric Univalence and Directed Motivation

2.1 Voevodsky’s univalence axiom

We recall the symmetric statement, both for orientation and to highlight what *breaks* when one passes to the directed setting.

Definition 2.1 (Equivalence). For $A, B : \mathcal{U}$, an *equivalence* is a function $f : A \rightarrow B$ together with a witness that the fibre $\sum_{x:A} f(x) =_B y$ is contractible for every $y : B$. We write $A \simeq B$ for the type of equivalences.

Definition 2.2 (idtoeqv). For each $A, B : \mathcal{U}$, the path induction principle defines a map

$$\text{idtoeqv}_{A,B} : (A =_{\mathcal{U}} B) \rightarrow (A \simeq B)$$

sending refl_A to the identity equivalence $\text{id}_A : A \simeq A$.

Theorem 2.3 (Voevodsky’s univalence axiom). *For every $A, B : \mathcal{U}$ the function $\text{idtoeqv}_{A,B}$ is itself an equivalence.*

The axiom yields its inverse $\text{ua} : (A \simeq B) \rightarrow (A =_{\mathcal{U}} B)$ and the computation $\text{idtoeqv} \circ \text{ua} \sim \text{id}$. Operationally, every equivalence delivers a path in the universe, and transport along this path acts as the equivalence.

What is symmetric. Two structural features of $A =_{\mathcal{U}} B$ make it manifestly *symmetric*:

1. It has *inverses*: if $p : A =_{\mathcal{U}} B$ then $p^{-1} : B =_{\mathcal{U}} A$.
2. It is *recursive*: idtoeqv may be defined by path induction, which is symmetric in A and B .

Both features are absent in the directed setting we wish to formalise.

2.2 What “directed” should mean

A directed type theory ought to admit, for each type A and pair of terms $a, b : A$, a hom-type

$$\text{Hom}_A(a, b) : \mathcal{U},$$

together with composition

$$\circ : \text{Hom}_A(b, c) \times \text{Hom}_A(a, b) \longrightarrow \text{Hom}_A(a, c)$$

and identities $\text{id}_a : \text{Hom}_A(a, a)$, satisfying associativity and unit laws *up to higher coherent homotopy*. Crucially, $\text{Hom}_A(a, b)$ should not in general carry inverses: if it did, the theory would collapse back to a symmetric one.

We can articulate desiderata.

(D1) *Hom is non-invertible.* The composition is not in general invertible.

(D2) *Functoriality.* Functions $f : A \rightarrow B$ act on hom-types:

$$f_* : \text{Hom}_A(a, b) \longrightarrow \text{Hom}_B(f(a), f(b)).$$

(D3) *Equivalences are identities.* For appropriately structured types, an invertible morphism in $\text{Hom}_A(a, b)$ should yield an identity $a =_A b$.

(D4) *Synthetic $(\infty, 1)$ -category theory*. The Yoneda lemma, adjoint functors, Kan extensions, and limits/colimits should be expressible *synthetically* on top of these primitives.

(D5) *Directed univalence*. The universe should reflect its own hom-structure: hom in the universe = function type.

(D1)–(D4) are addressed by Riehl–Shulman and the body of subsequent work [3, 5, 7, 9]. (D5) is the topic of this paper.

2.3 $(\infty, 1)$ -categorical motivation

There is a clean semantic picture explaining why directed univalence is hard. In HoTT, the universe $\mathcal{U}$ is interpreted as an object classifier in an $(\infty, 1)$ -topos: $\mathcal{U}$ classifies small bundles up to fibrewise equivalence. Univalence is the corresponding internal statement: equivalences of bundles are paths in $\mathcal{U}$.

The directed analogue would have a universe $\mathcal{S}$ classifying bundles *up to map*, not just equivalence. In an $(\infty, 1)$ -topos $\mathcal{X}$, the natural such classifier is the cocartesian fibration of the $(\infty, 1)$ -category $\mathcal{X}_{/-}$ over $\mathcal{X}$ itself. But this is no longer a single object: it is a structured presentable $(\infty, 1)$ -category. The naive analogue — “a small object in $\mathcal{X}$ that classifies maps” — requires *at least* a handle on $\mathcal{X}$ as an $(\infty, 1)$ -category internal to $\mathcal{X}$, a strictly stronger demand than internal ∞ -groupoids.

This is the structural reason a complete directed univalence principle is hard: it forces the type theory to internalise its own large $(\infty, 1)$ -categorical structure, not merely its ∞ -groupoidal one.

3 Riehl–Shulman Simplicial Type Theory

3.1 Topes, cubes, extension types

STT is a three-layer system: a homotopy type theory; a strict cube layer with the directed interval $\mathbf{2} = \{0 \leq 1\}$; and a tope layer of decidable propositions over the cube. We sketch each layer.

Definition 3.1 (Cube). A *cube* is built inductively from a primitive directed interval $\mathbf{2}$, the unit cube $\mathbf{1}$, and finite products. Cubes are denoted I, J, K .

Definition 3.2 (Tope). A *tope* on a cube I is a polynomial inequality in the coordinates of I , closed under $\top, \perp, \wedge, \vee$. Topes form a distributive lattice. The standard examples on $\mathbf{2}$ are $\{0\}$, $\{1\}$, $\{0 \leq t\}$, $\{t \leq 1\}$.

Definition 3.3 (Extension type). Given a cube I , a tope Φ on I , a type family $A : I \rightarrow \mathcal{U}$, and a partial term $a : \prod_{t:\Phi} A(t)$ defined on Φ , the *extension type*

$$\langle \prod_{t:I} A(t) \mid \Phi \mapsto a \rangle$$

is the type of dependent functions $f : \prod_{t:I} A(t)$ such that for $t : \Phi$, $f(t) \equiv a(t)$ judgmentally.

Remark 3.4. Extension types provide a uniform mechanism for *strict boundary conditions*: the value of f on the tope Φ is fixed up to definitional equality. This is what lets STT speak about, e.g., morphisms with prescribed source and target.

3.2 Hom-types and Segal types

Definition 3.5 (Hom-type). Let $A : \mathcal{U}$ and $a, b : A$. The *hom-type* is the extension type

$$\mathrm{Hom}_A(a, b) := \left\langle \prod_{t:2} A \mid \{0\} \mapsto a, \{1\} \mapsto b \right\rangle.$$

So a morphism from a to b is a function $2 \rightarrow A$ taking 0 to a and 1 to b . The directed interval is *ordered*, hence so is the hom-type.

Definition 3.6 (Segal predicate). A type A is *Segal*, written $\mathrm{isSegal}(A)$, if the canonical map

$$\mathrm{Hom}_A(a, b) \times \mathrm{Hom}_A(b, c) \rightarrow \left\langle \prod_{(t,s):\Lambda_1^2} A \mid (a, \cdot, c) \right\rangle$$

into the type of Λ_1^2 -fillers is an equivalence, where Λ_1^2 is the inner horn of the standard 2-simplex — concretely, the sub-cube of the unit square 2×2 given by the tope $\{(t, s) \mid t \leq s\}$ minus its long edge from $(0, 0)$ to $(1, 1)$, equivalently the shape generated by the two boundary edges $(t, 0)$ and $(1, s)$ [3, §5]. In particular, every pair of composable morphisms admits a unique-up-to-coherent-homotopy composite.

Segal types are the STT version of Segal complete spaces. Composition is then a definable operation:

$$\circ : \mathrm{Hom}_A(b, c) \times \mathrm{Hom}_A(a, b) \rightarrow \mathrm{Hom}_A(a, c)$$

when $\mathrm{isSegal}(A)$.

Definition 3.7 (Rezk predicate). A Segal type A is *Rezk* if the canonical map

$$(a =_A b) \rightarrow \mathrm{Iso}_A(a, b)$$

is an equivalence, where $\mathrm{Iso}_A(a, b) \subset \mathrm{Hom}_A(a, b)$ is the sub-type of two-sidedly invertible morphisms. (This is internal univalence on A .)

Theorem 3.8 (Riehl–Shulman, semantic). *Rezk types in **STT** correspond, in the semantic interpretation in bisimplicial sets, to complete Segal spaces, and hence to $(\infty, 1)$ -categories.*

3.3 Discrete types

Definition 3.9 (Discrete type). A type A is *discrete*, written $\mathrm{isDiscrete}(A)$, if for every $a, b : A$ the canonical map

$$(a =_A b) \rightarrow \mathrm{Hom}_A(a, b)$$

is an equivalence; equivalently, every morphism in A is invertible. Discrete types are the STT counterpart of ∞ -groupoids.

Example 3.10. The natural numbers $\mathbb{N}$ are discrete in **STT** (they carry no non-trivial 1-cells). More generally, every type imported from a model of HoTT into **STT** via the discrete inclusion is discrete.

3.4 Synthetic Yoneda lemma

A high point of [3] is a synthetic Yoneda lemma. We state it informally.

Theorem 3.11 (Synthetic Yoneda, RS Thm 9.1). *Let A be a Segal type and $\mathcal{S}$ the universe of discrete types. Then for any covariant functor $F : A \rightarrow \mathcal{S}$ and any $a : A$, evaluation at id_a gives an equivalence*

$$\left(\prod_{x:A} \text{Hom}_A(a, x) \rightarrow F(x) \right) \simeq F(a).$$

Proof sketch. Both sides are recovered from the dependent product structure on A and the universal property of the representable $\text{Hom}_A(a, -)$. Detailed proof: see [3, §9]. $\square$

This is the engine that drives much of synthetic $(\infty, 1)$ -category theory in **STT**.

4 Partial Directed Univalence (Riehl–Shulman 2017)

4.1 Covariant fibrations

Definition 4.1 (Covariant fibration). Let $A : \mathcal{U}$ be a Segal type. A type family $P : A \rightarrow \mathcal{U}$ is *covariant* if for each $\alpha : \text{Hom}_A(a, b)$ and each $u : P(a)$ there is a chosen lift $\alpha_*(u) : P(b)$, depending functorially on α .

In semantic terms, covariant families correspond to *left fibrations*: functors $A \rightarrow \mathcal{S}$ classified by a particular kind of straightening.

Theorem 4.2 (Straightening, partial directed univalence for covariant families). *For a Rezk type A , the following types are equivalent:*

1. *Covariant type families $P : A \rightarrow \mathcal{U}$ valued in discrete types.*
2. *Functors $A \rightarrow \mathcal{S}_{\text{disc}}$, where $\mathcal{S}_{\text{disc}}$ is the universe of discrete types.*

This is *directed univalence for covariant fibrations valued in discrete types*: the universe $\mathcal{S}_{\text{disc}}$ classifies covariant families, in exact analogy with $\mathcal{U}$ classifying type families up to equivalence.

4.2 What is partial about it

Theorem 4.2 stops short of asserting

$$\text{Hom}_{\mathcal{S}}(A, B) \simeq \text{Fun}(A, B)$$

for the universe $\mathcal{S}$ *itself* treated as a Segal type. Two reasons.

(1) The universe is not internally Segal. In **STT** as set up by Riehl–Shulman, the discrete universe $\mathcal{S}_{\text{disc}}$ is given as a sub-type of $\mathcal{U}$, but its hom-type $\text{Hom}_{\mathcal{S}_{\text{disc}}}(A, B)$ is not directly described. The straightening theorem describes covariant *maps into* $\mathcal{S}_{\text{disc}}$, not the hom-type *within* $\mathcal{S}_{\text{disc}}$.

(2) Modal pieces are missing. To formulate “ $\mathrm{Hom}_{\mathcal{S}}(A, B) \simeq \mathrm{Fun}(A, B)$ ” coherently, one needs to mediate between the directed setting (where the universe carries non-trivial morphisms) and the symmetric setting (where it carries only equivalences). This is precisely the function of the modal apparatus of Gratzer–Weinberger–Buchholtz, which is absent from [3].

4.3 Limitations made formal

Proposition 4.3 (Limitations of **STT**). *In Riehl–Shulman simplicial type theory:*

1. *There is a directed straightening for covariant families into the discrete universe (Theorem 4.2).*
2. *There is no internally definable hom-type for the discrete universe $\mathcal{S}_{\mathrm{disc}}$ qua object of the universe $\mathcal{U}$ that satisfies $\mathrm{Hom}_{\mathcal{S}_{\mathrm{disc}}}(A, B) \simeq \mathrm{Fun}(A, B)$.*
3. *There is no Segal-completion theorem for $\mathcal{S}_{\mathrm{disc}}$ producing a Segal-type version with the desired hom.*

These limitations motivate the modal extension we now describe.

5 Triangulated Type Theory (Gratzer–Weinberger–Buchholtz 2024)

5.1 Modal layer

The crucial new ingredient in [4] is a system of modalities. We describe four.

Definition 5.1 (The four modalities). **TTT** extends **STT** with modalities:

1. $\flat$ “crisp” / “discrete-discrete”: forces a type to be discrete and only depends on globally defined data.
2. $\sharp$ “codiscrete”: right adjoint to $\flat$.
3. **Disc**: the modality whose modal types are the discrete types.
4. $(-)^{\mathrm{op}}$: the opposite modality reversing direction.

With $\flat \dashv \mathrm{Disc} \dashv \sharp$ a triangle of cohesion-like modalities, and $(-)^{\mathrm{op}}$ as duality.

Remark 5.2. The $\flat \dashv \mathrm{Disc} \dashv \sharp$ adjunction is structurally identical to the cohesive triple in Schreiber–Shulman cohesion. The new aspect is the interaction with $(-)^{\mathrm{op}}$, which has no analogue in cohesive HoTT.

5.2 The discrete universe $\mathcal{S}$

Definition 5.3 (Discrete universe). The *discrete universe* $\mathcal{S}$ in **TTT** is the type of Disc-modal types: $\mathcal{S} := \sum_{A:\mathcal{U}} \text{isDiscrete}(A)$, with universe structure inherited from $\mathcal{U}$.

Theorem 5.4 (GWB, $\mathcal{S}$ is Segal). *The discrete universe $\mathcal{S}$ is a Segal type in **TTT**. The hom-type $\text{Hom}_{\mathcal{S}}(A, B)$ is well-defined as an extension type out of $\mathbf{2}$ valued in $\mathcal{S}$.*

Proof outline. The proof uses the modal layer to control how the family $A : \mathbf{2} \rightarrow \mathcal{S}$ behaves at the endpoints. Key step: the dependent product over $\mathbf{2}$ in $\mathcal{S}$ exists because the discreteness predicate is preserved by extension types, which is established via the $\flat$ modality. See [4, §6]. $\square$

5.3 The directed univalence theorem

Theorem 5.5 (Directed univalence in **TTT**, GWB Thm 8.1). *For all $A, B : \mathcal{S}$, the canonical map*

$$\text{funtohom}_{A,B} : \text{Fun}(A, B) \longrightarrow \text{Hom}_{\mathcal{S}}(A, B)$$

is an equivalence. Concretely, every directed morphism in the universe of discrete types is uniquely (up to coherent homotopy) the type-theoretic action of a function.

Proof sketch. The proof has three parts.

1. *Define funtohom.* Given $f : A \rightarrow B$, define the family $A_f : \mathbf{2} \rightarrow \mathcal{S}$ by $A_f(0) := A$, $A_f(1) := B$, $A_f(t) := \text{cofib}(f, t)$, where $\text{cofib}(f, t)$ denotes a standard cofibrant path-object construction interpolating between A and B along f (concretely, the homotopy pushout of f with the trivial map at parameter t). The Disc modality ensures $A_f(t)$ stays discrete for $t \in (0, 1)$.
2. *Inverse construction.* Given $\alpha : \text{Hom}_{\mathcal{S}}(A, B)$, extract the function $\alpha^{\flat} : A \rightarrow B$ via the $\flat$ modality applied at the endpoint $\{1\}$.
3. *Coherence.* Composition $\text{funtohom} \circ \flat = \text{id}$ and $\flat \circ \text{funtohom} = \text{id}$ both follow from the universal property of extension types and the Beck–Chevalley conditions for $\flat \dashv \text{Disc} \dashv \sharp$.

See [4, §8] for full details. $\square$

Corollary 5.6 (Yoneda for $\mathcal{S}$). *Every covariant functor $F : \mathcal{S} \rightarrow \mathcal{S}$ is naturally equivalent to $\text{Hom}_{\mathcal{S}}(F^{\flat}(\mathbf{1}), -)$ for a uniquely determined object $F^{\flat}(\mathbf{1}) : \mathcal{S}$.*

Remark 5.7. Theorem 5.5 is a complete directed univalence statement for the discrete fragment. It is “triangulated” in the sense that its formulation requires the triangle $\flat \dashv \text{Disc} \dashv \sharp$.

5.4 Worked examples in the discrete universe

We illustrate Theorem 5.5 with three small examples that occupy the discrete universe $\mathcal{S}$.

Example 5.8 (Hom in $\mathcal{S}$ between finite types). Let $A := \mathbf{2} = \{0, 1\}$ and $B := \mathbf{3} = \{0, 1, 2\}$, both discrete. Then $\text{Fun}(A, B) = B^A$ has 9 elements. Theorem 5.5 asserts $\text{Hom}_{\mathcal{S}}(A, B) \simeq \text{Fun}(A, B)$, so the hom-type from A to B in $\mathcal{S}$ has, up to homotopy, exactly nine path-components, each indexed by a function $\mathbf{2} \rightarrow \mathbf{3}$.

Example 5.9 (Hom from $\mathbb{N}$ to $\mathbb{N}$ in $\mathcal{S}$). Let $A = B = \mathbb{N}$. Then $\text{Hom}_{\mathcal{S}}(\mathbb{N}, \mathbb{N}) \simeq \mathbb{N} \rightarrow \mathbb{N}$, an uncountable type. Each function $f : \mathbb{N} \rightarrow \mathbb{N}$ corresponds to a unique directed morphism, e.g. the successor function gives a canonical morphism $\text{succ}_* : \mathbb{N} \rightarrow \mathbb{N}$ in $\mathcal{S}$.

Example 5.10 (Composition agrees with composition). Functoriality of `funtohom` ensures $\text{funtohom}_{A,C}(g \circ f) = \text{funtohom}_{B,C}(g) \circ \text{funtohom}_{A,B}(f)$. So composition in $\mathcal{S}$ models composition of functions exactly. This is the property that earns $\mathcal{S}$ the name “the synthetic universe of discrete types”.

Remark 5.11 (Naturality). The map `funtohom` is natural in both arguments: pre- and post-composition with a function $f : A' \rightarrow A$, $g : B \rightarrow B'$ commute with the equivalences. This is what allows $\mathcal{S}$ to be regarded as an internally-defined Segal type whose morphisms reproduce the functions of the ambient HoTT.

5.5 Comparison to the symmetric universe

It is instructive to compare Theorem 5.5 with Voevodsky’s univalence (Theorem 2.3).

- *Voevodsky*: $\text{Id}_{\mathcal{U}}(A, B) \simeq (A \simeq B)$. Identifications in the universe correspond to equivalences of types.
- *GWB*: $\text{Hom}_{\mathcal{S}}(A, B) \simeq \text{Fun}(A, B)$. Directed morphisms in the universe of discrete types correspond to functions of types.

In a sense, GWB is structurally simpler: $\text{Fun}(A, B)$ is the standard function type, while $A \simeq B$ requires the entire equivalence apparatus. But GWB only applies to the *discrete* sub-universe, so the principle is local. The challenge of extending it (Section 6) is to find a richer universe in which non-equivalence functions still correspond to directed morphisms.

5.6 Implementation status in `rzk`

The `rzk` proof assistant [11] implements simplicial type theory with topes and extension types. It does *not* (as of 2026) implement triangulated modalities or prove directed univalence. Bubenik et al. are pursuing a triangulated extension; partial formalisations of the GWB theorem are in progress.

The Haskell module `Hom.hs` accompanying this paper exhibits a toy form of `funtohom` for the discrete universe in a Hask-like setting. The Lean 4 file `Directed.lean` postulates Theorem 5.5 as an axiom and derives small consequences.

5.7 What is and is not proved

Proposition 5.12 (Scope of GWB). *Theorem 5.5 establishes directed univalence for the discrete universe $\mathcal{S}$. It does not establish:*

1. *Directed univalence for any non-discrete sub-universe of $\mathcal{U}$.*
2. *A directed analogue of Voevodsky univalence for the universe $\mathcal{U}$ itself.*
3. *A universe of $(\infty, 1)$ -categories (Rezk types) with hom-types behaving as functors.*

6 Toward a Complete Principle

We now formulate four increasingly strong candidate principles and analyse what each proof would require.

6.1 Four candidates

Definition 6.1 (Candidate I: Discrete directed univalence). *Candidate I.* For the discrete universe $\mathcal{S}$, $\text{Hom}_{\mathcal{S}}(A, B) \simeq \text{Fun}(A, B)$.

Status: Proven (Theorem 5.5, GWB 2024).

Definition 6.2 (Candidate II: Segal directed univalence). *Candidate II.* For the universe $\mathcal{S}_{\text{Seg}}$ of Segal types, $\text{Hom}_{\mathcal{S}_{\text{Seg}}}(A, B) \simeq \text{SegFun}(A, B)$, where SegFun is the type of functors preserving Segal structure (= ordinary functions, since composition is uniquely determined).

Status: Open. Conjectured.

Definition 6.3 (Candidate III: Rezk directed univalence). *Candidate III.* For the universe $\mathcal{S}_{\text{Rezk}}$ of Rezk types, $\text{Hom}_{\mathcal{S}_{\text{Rezk}}}(A, B) \simeq \text{Fun}(A, B)$.

Status: Open. The principal target.

Definition 6.4 (Candidate IV: Universal directed univalence). *Candidate IV.* For the entire universe $\mathcal{U}$, equipped with the directed structure inherited from extension types out of **2**, $\text{Hom}_{\mathcal{U}}(A, B) \simeq \text{Fun}(A, B)$.

Status: Open. Likely requires a 2-level theory.

6.2 Layered programme

Conjecture 6.5 (Layered directed univalence). There is a tower of universes

$$\mathcal{S}_{\text{disc}} \subset \mathcal{S}_{\text{Rezk}} \subset \mathcal{S}_{\text{Seg}} \subset \mathcal{U}$$

such that each level satisfies a directed univalence principle of its own kind, with the inclusions being conservative.

We articulate a layered *programme* as follows. Let L_n denote a universe at layer n . Then:

- $L_0 = \mathcal{S}_{\text{disc}}$: classical directed univalence (Candidate I).

- $L_1 = \mathcal{S}_{\text{Seg}}$: Segal directed univalence (Candidate II); morphisms are Segal functors.
- $L_2 = \mathcal{S}_{\text{Rezk}}$: Rezk directed univalence (Candidate III); morphisms are functors of $(\infty, 1)$ -categories.
- $L_3 = \mathcal{U}$: universal directed univalence (Candidate IV); requires reflective sub-universe technology.

6.3 Obstructions to Candidates II–IV

Proposition 6.6 (Obstruction: opposite-type coherence). *Candidates II–IV require a coherent theory of opposite types A^{op} . The opposite type construction in **STT** is well-defined, but its compatibility with the universe structure is subtle: $(-)^{\text{op}}$ is a non-trivial involution on $\mathcal{U}$ that interacts non-trivially with Hom .*

Proposition 6.7 (Obstruction: modal dependence). *Candidates II–IV likely require a modal layer richer than $\flat \dashv \text{Disc} \dashv \sharp$. In particular, the modality picking out Rezk types is not in general a left or right adjoint to either $\flat$ or Disc , so cannot be expressed in cohesion-like form.*

Proposition 6.8 (Obstruction: large vs. small). *Candidate IV in its naive form is inconsistent: $\mathcal{U}$ cannot internalise its own “large” $(\infty, 1)$ -categorical structure. A 2-level approach (à la Annenkov–Capriotti–Kraus–Sattler) appears necessary, with a strict inner universe and a homotopical outer one.*

6.4 Sketch of a proof attempt for Candidate III

We sketch what a proof of Candidate III would look like, identifying the gaps.

1. Define the universe $\mathcal{S}_{\text{Rezk}}$ as the sub-type of $\mathcal{U}$ on Rezk types: $\mathcal{S}_{\text{Rezk}} := \sum_{A:\mathcal{U}} \text{isRezk}(A)$.
2. Show $\mathcal{S}_{\text{Rezk}}$ is Segal in **TTT**. This is plausibly an extension of Theorem 5.4 but requires that isRezk commute with extension types out of **2**, which is non-trivial.
3. Show $\mathcal{S}_{\text{Rezk}}$ is in fact Rezk (so that “equivalences in $\mathcal{S}_{\text{Rezk}}$ are identities”). Equivalences in this context are functors that are essentially surjective and fully faithful in the synthetic sense.
4. Construct $\text{funtohom} : \text{Fun}(A, B) \rightarrow \text{Hom}_{\mathcal{S}_{\text{Rezk}}}(A, B)$. Functions between Rezk types are functors of $(\infty, 1)$ -categories.
5. Construct an inverse using a twisted modality $\text{Disc}^{\text{Rezk}}$ that picks out Rezk types as a localisation of **STT**.

Step 5 is the principal technical gap. The modality $\text{Disc}^{\text{Rezk}}$ does not exist in published **TTT**.

Question 6.9. Does there exist an extension of **TTT** with a modality $\text{Disc}^{\text{Rezk}}$ such that:

1. Modal types are exactly Rezk types.
2. The triangle $\flat \dashv \text{Disc}^{\text{Rezk}} \dashv \sharp_{\text{Rezk}}$ extends the discrete cohesion.
3. funtohom is an equivalence?

6.5 The Cisinski–Nguyen line of work

A complementary semantic line of work is due to Cisinski and Nguyen [8], who construct a *universal cocartesian fibration* for the $(\infty, 1)$ -topos of presheaves of spaces. In categorical terms, this is the directed analogue of an object classifier: it is a cocartesian fibration $\mathcal{U}^\rightarrow \rightarrow \mathcal{U}$ whose fibre over $A : \mathcal{U}$ is essentially the slice $\mathcal{U}_{/A}$.

Translated into the directed type theory programme, Cisinski–Nguyen suggests a *semantic* sub-universe of $\mathcal{U}$ on which the directed-univalence statement should hold:

$$\mathrm{Hom}_{\mathcal{U}^\rightarrow}(A, B) \simeq \mathcal{U}(A, B) = \mathrm{Fun}(A, B).$$

The question is whether this sub-universe can be characterised *syntactically* in a directed type theory, and in particular whether the modal apparatus of **TTT** extends to it.

Conjecture 6.10 (Conjectural Cisinski–Nguyen internal language). There exists an extension of **TTT** in which the universal cocartesian fibration is internally definable, and Candidate III (Theorem 6.3) holds for the resulting universe of types.

Remark 6.11. Even if Conjecture 6.10 is true, proving it likely requires non-trivial extensions to the modal layer: at minimum, modalities for the slice $\mathcal{U}_{/A}$ uniformly in A are needed.

6.6 Cubical comparison

Remark 6.12 (Cubical analogue). Voevodsky univalence is a postulate in Book HoTT but a *theorem* in cubical type theory (CCHM, Cohen–Coquand–Huber–Mörtberg) due to Glue types. The analogous question for the directed setting is open: is there a cubical-flavoured presentation of **TTT** (or its successors) in which directed univalence is constructive? Cisinski has suggested fragments [7]; the rzk proof assistant aims at a computational interpretation [11]. The Weaver–Licata bicubical-set model [9] provides a constructive directed univalence for a particular fragment of bicubical type theory, indicating that a fully computational directed univalence is achievable in restricted settings.

6.7 An alternative axiomatic approach

Rather than seeking a syntactic-by-construction proof, one can take Theorem 5.5 as a model and *postulate* a version of directed univalence at each layer. Concretely:

1. Add to **TTT** a constant $\mathrm{dua}_{\mathrm{Rezk}} : \prod_{A, B : \mathcal{S}_{\mathrm{Rezk}}} \mathrm{Fun}(A, B) \rightarrow \mathrm{Hom}_{\mathcal{S}_{\mathrm{Rezk}}}(A, B)$.
2. Add an axiom that $\mathrm{dua}_{\mathrm{Rezk}}$ is an equivalence.
3. Verify in standard models (bisimplicial sets, marked simplicial sets, ∞ -toposes) that the postulated principle is sound.

This is the directed analogue of taking Voevodsky’s univalence as an axiom rather than deriving it from cubical methods. It establishes consistency with prevailing semantic models but leaves the syntactic story incomplete.

Proposition 6.13 (Conjectural soundness of axiomatic directed univalence). *The axiomatic approach above is sound in any $(\infty, 1)$ -topos $\mathcal{X}$ that admits a small object classifier and is bicartesian.*

Proof outline. Such an $\mathcal{X}$ admits a sub- $(\infty, 1)$ -category of internal Rezk objects. The hom-spaces in this sub-category are by construction the function-types of the underlying ∞ -topos. Soundness then follows from the universality of the construction. $\square$

6.8 Summary table

Universe	Variant	Status
$\mathcal{U}$	symmetric	Voevodsky univalence (theorem in cubical)
$\mathcal{S}_{\text{disc}}$	directed, discrete	GWB 2024 (Theorem 5.5)
$\mathcal{S}_{\text{Seg}}$	directed, Segal	Open (Cand. II)
$\mathcal{S}_{\text{Rezk}}$	directed, Rezk	Open (Cand. III, principal)
$\mathcal{U}_{\text{dir}}$	directed, all types	Open (Cand. IV, requires 2-level)

7 Connections to $(\infty, 1)$ -Natural Number Objects and HITs

7.1 Natural number objects in directed type theory

In ordinary HoTT, the natural numbers $\mathbb{N}$ are characterised by an initiality universal property: $(\mathbb{N}, 0, \text{succ})$ is initial in the type of pointed endo-types. This characterisation generalises directly to the discrete sub-universe of **TTT**: $\mathbb{N}$ is a type in $\mathcal{S}$ and is initial.

Proposition 7.1 (NNO in discrete directed type theory). *In **TTT**, the natural numbers $\mathbb{N} : \mathcal{S}$ are characterised up to equivalence by initiality:*

$$\text{isContr} \left(\sum_{f: \mathbb{N} \rightarrow A} f(0) =_A a_0, \text{ compatible with succ} \right)$$

for any pointed endo-discrete-type (A, a_0, s) . Directed univalence (Theorem 5.5) then identifies these initiality witnesses with directed morphisms in $\mathcal{S}$.

7.2 Higher inductive types and directed univalence

A higher inductive type (HIT) like S^1 has constructors and *path constructors* producing identifications. In the directed setting one wants *directed-path constructors*: a type with constructors and prescribed hom-types, defined synthetically.

Definition 7.2 (Directed HIT). A *directed higher inductive type* is a type generated by point constructors plus directed-edge constructors valued in **Hom**-types.

Example 7.3 (Walking arrow). The simplest directed HIT is the *walking arrow* **2** itself: two points $0, 1$ and a generator $\sigma : \text{Hom}_2(0, 1)$.

Example 7.4 (Walking commutative square). Four points 00, 01, 10, 11, four directed edges, and a 2-cell witnessing commutativity. Models the $\Delta^1 \times \Delta^1$ shape.

7.3 Synthetic representation theory

Directed univalence connects to representation theory through the slogan: *representations are functors*. A representation of a group G on a vector space is a functor $BG \rightarrow \mathbf{Vect}$, where BG is the classifying type of G . In a hypothetical full directed type theory:

Definition 7.5 (Synthetic representation). A *synthetic representation* of a Rezk-type group G is a directed morphism $BG \rightarrow \mathbf{Vect}$ in $\mathcal{S}_{\text{Rezk}}$. Here, as in classical higher category theory, the group G is presented through its *delooping* BG — the one-object Rezk-type whose endomorphism hom-type is G itself; representations are then functors out of this delooping. Sub-representations are directed monomorphisms in this sense.

Remark 7.6. Eventually, automorphic representations (functors $\text{GL}(n, \mathbb{A}) \rightarrow \infty\text{-Grpd}$ for $\mathbb{A}$ adelic) would be morphisms in a directed universe. This is highly speculative and depends on Candidate III.

7.4 Synthetic adjunctions and Kan extensions

Directed univalence enables a direct treatment of adjunctions in synthetic $(\infty, 1)$ -category theory. In the GWB setting, an adjunction $F \dashv G : A \rightarrow B$ between objects of the directed universe of Rezk types becomes a single piece of data in $\mathcal{S}_{\text{Rezk}}$:

Definition 7.7 (Synthetic adjunction). An *adjunction* between Rezk-types $A, B : \mathcal{S}_{\text{Rezk}}$ consists of a pair $F : A \rightarrow B$, $G : B \rightarrow A$ together with directed natural transformations

$$\eta : \text{id}_A \Rightarrow G \circ F, \quad \varepsilon : F \circ G \Rightarrow \text{id}_B$$

satisfying the standard triangle identities. By Candidate III, the data is equivalent to a single object in a derived directed universe of adjunctions.

Remark 7.8. Right Kan extension $\text{Ran}_p F$ along $p : C \rightarrow D$ would similarly be a directed morphism in $\mathcal{S}_{\text{Rezk}}$ rather than a global postulate. The whole package of Mac Lane–Lurie $(\infty, 1)$ -category theory becomes type-theoretically internal.

7.5 Stability under base change

A subtle but important property of complete directed univalence is its behaviour under base change. If $f : C \rightarrow D$ is a directed morphism in $\mathcal{S}_{\text{Rezk}}$, the induced functor $f^* : \mathcal{S}_{\text{Rezk}}/D \rightarrow \mathcal{S}_{\text{Rezk}}/C$ should preserve directed-univalent structure.

Proposition 7.9 (Base-change stability, conjectural). *Assume Candidate III. Then for every directed morphism $f : C \rightarrow D$ in $\mathcal{S}_{\text{Rezk}}$, the pullback f^* commutes with funtohom up to coherent equivalence.*

Proof sketch. Pullback in $\mathcal{S}_{\text{Rezk}}$ corresponds semantically to pullback of cocartesian fibrations, which is functorial. The functoriality lifts funtohom to a natural equivalence between fibres. $\square$

7.6 Comparison to formalisations: Coq-HoTT, Lean-mathlib, Cubical Agda, rzk

- Coq-HoTT and Lean-mathlib formalise symmetric univalence as a postulate.
- Cubical Agda makes symmetric univalence *computational*.
- The experimental proof assistant rzk [11] implements **STT** with toposes and extension types; it does *not* yet implement triangulated modalities or prove directed univalence.

8 Discussion

8.1 Why is directed univalence so hard?

We can summarise the difficulty in three points.

1. *Asymmetry of hom.* The hom-type $\mathbf{Hom}_A(a, b)$ is asymmetric in (a, b) , while the identity type $a =_A b$ is symmetric. Many path-induction-based arguments fail directly.
2. *Modal dependence.* Even the discrete fragment requires modalities to mediate between symmetric and directed worlds. Non-discrete fragments will likely need richer modal apparatus.
3. *Size issue.* The universe is a large object. A directed structure on it would package its $(\infty, 1)$ -categorical morphisms internally — a strictly stronger demand than packaging ∞ -groupoidal equivalences.

8.2 What would change if Candidate III were proven?

1. Synthetic $(\infty, 1)$ -category theory in **STT/TTT** would gain a *universe-theoretic* foundation, not just a Yoneda-based one.
2. Functors $A \rightarrow B$ between Rezk types would automatically be morphisms in a Rezk universe, enabling direct manipulation of categories of categories.
3. Adjunctions would internalise: an adjunction $F \dashv G$ would correspond to a single object in $\mathbf{Adj}(\mathcal{S}_{\mathbf{Rezk}})$, accessible through directed transport.
4. Formalisation in a directed proof assistant (extension of rzk) would become tractable.

8.3 What would change if Candidate IV were proven?

Candidate IV is qualitatively stronger. If proven (likely in a 2-level theory):

1. The universe itself would be both an ∞ -groupoid and an $(\infty, 1)$ -category — with the structures coherently related.
2. Univalence and directed univalence would be compatible.
3. Higher directed univalence (for (∞, n) -categories) could be tackled inductively.

8.4 Limitations of this paper

This paper is expository and programmatic. We do not prove a new directed univalence theorem. The accompanying Haskell and Lean code formalise the discrete fragment of **TTT** in a toy model (Haskell) and as a postulate-based sketch (Lean), and exhibit the algebraic content. We hope these provide a productive substrate for future work.

9 Open Problems

We list four explicit open problems.

Question 9.1 (Rezk directed univalence). Prove or disprove Candidate III: $\mathrm{Hom}_{\mathcal{S}_{\mathrm{Rezk}}}(A, B) \simeq \mathrm{Fun}(A, B)$ for all Rezk types A, B .

Question 9.2 ($(\infty, 1)$ -Cat as a primitive type). Is there an extension of **TTT** with a primitive type $\mathbf{Cat}_{\infty, 1}$ of $(\infty, 1)$ -categories, satisfying its own directed univalence?

Question 9.3 (Internal language of cocartesian fibrations). Develop the internal language of cocartesian fibrations in **TTT**, in the spirit of Cisinski–Nguyen [8].

Question 9.4 (Directed Langlands). Reformulate automorphic representations as directed morphisms $\mathrm{GL}(n, \mathbb{A}) \rightarrow \infty\text{-Grpd}$ in a fully directed type theory.

10 Conclusion

Directed univalence has come a long way since Riehl–Shulman 2017. Their original framework gave us simplicial type theory with an internalised hom-type and a partial straightening theorem for covariant fibrations. Gratzer–Weinberger–Buchholtz 2024 introduced the modal apparatus needed to internalise directed structure on the universe, and proved the first complete directed univalence theorem for the universe of discrete types.

The complete principle remains a research target. The path forward is layered: extend the modality structure to capture Segal-ness and Rezk-ness, prove successively stronger directed univalence statements at each layer, and ultimately reach a 2-level account where the universe carries both directed and undirected structure coherently. The reward, if it can be obtained, is a synthetic foundation for $(\infty, 1)$ -category theory in which functors are directed morphisms in the universe — an *ipso facto* foundation for synthetic representation theory and beyond.

We close with a final observation. The classical univalence axiom internalised the slogan “isomorphic objects are equal”. A complete directed univalence axiom would internalise the slogan “functors are arrows in the universe of categories”. These are both deeply structuralist principles. The latter, if it can be proved, would unify the structuralist viewpoints of Lawvere, Awodey, and Voevodsky into a single foundational system in which the universe is itself a category.

Acknowledgments

This is a derivative exposition; original results are due to Emily Riehl, Michael Shulman, Daniel Gratzer, Jonathan Weinberger, and Ulrik Buchholtz. We thank these authors for the foundational work surveyed here.

References

- [1] The Univalent Foundations Program. *Homotopy Type Theory: Univalent Foundations of Mathematics*. Institute for Advanced Study, 2013.
- [2] V. Voevodsky. *The univalence axiom*. Talks and lecture notes, 2010–2014.
- [3] E. Riehl and M. Shulman. “A type theory for synthetic ∞ -categories.” *Higher Structures* 1(1):147–224, 2017. arXiv:1705.07442.
- [4] D. Gratzer, J. Weinberger, and U. Buchholtz. “Directed univalence in simplicial homotopy type theory.” arXiv:2407.09146, 2024.
- [5] U. Buchholtz. *Higher structures in homotopy type theory*. Habilitation, TU Darmstadt, 2019.
- [6] E. Riehl. “Synthetic perspectives on spaces and categories.” Survey article, in preparation (2025); preprint announced on the author’s webpage.
- [7] D.-C. Cisinski. *Higher categories and homotopical algebra*. Cambridge Studies in Advanced Mathematics 180, 2019.
- [8] D.-C. Cisinski and H.-K. Nguyen. “The universal cocartesian fibration.” arXiv:2210.08945, 2022.
- [9] G. A. Kavvos and M. Weaver. “On directed type theory.” Ongoing work, 2024–2025; cf. M. Weaver and D. Licata, “A Constructive Model of Directed Univalence in Bicubical Sets,” *LICS 2020*.
- [10] M. Shulman. “All $(\infty, 1)$ -toposes have strict univalent universes.” arXiv:1904.07004, 2019.
- [11] N. Bubenik et al. *rzk: an experimental proof assistant for synthetic ∞ -category theory*. <https://rzk-lang.github.io/rzk/>, accessed 2026.
- [12] J. Lurie. *Higher Topos Theory*. Annals of Mathematics Studies 170, Princeton, 2009.
- [13] D. Annenkov, P. Capriotti, N. Kraus, and C. Sattler. “Two-level type theory and applications.” *Math. Struct. Comput. Sci.* 33(8), 2023.
- [14] P. Lumsdaine and M. Shulman. “Semantics of higher inductive types.” arXiv:1705.07088, 2017.

- [15] D. R. Licata and M. Shulman. “Calculating the fundamental group of the circle in homotopy type theory.” *LICS 2013*, pp. 223–232.
- [16] A. Joyal. “Quasi-categories and Kan complexes.” *J. Pure Appl. Algebra* 175(1):207–222, 2002.
- [17] C. Rezk. “A model for the homotopy theory of homotopy theory.” *Trans. AMS* 353(3):973–1007, 2001.
- [18] P. R. North. “Towards a directed homotopy type theory.” *ENTCS* 347, 2019.
- [19] S. Awodey. “Structuralism, invariance, and univalence.” *Phil. Math.* 22(1), 2014.