

# A logical framework for data-driven reasoning

PAOLO BALDI\*, *Department of Human Studies, University of Salento, Via Valesio, 24 73100 Lecce, Italy.*

ESTHER ANNA CORSI\*\*, *LUCI Lab, Department of Philosophy, University of Milan, Via Festa del Perdono, 7 20122 Milano, Italy.*

HYKEL HOSNI†, *LUCI Lab, Department of Philosophy, University of Milan, Via Festa del Perdono, 7 20122 Milano, Italy.*

## Abstract

We introduce and investigate a family of consequence relations with the goal of capturing certain important patterns of data-driven inference. The inspiring idea for our framework is the fact that data may reject, possibly to some degree, and possibly by mistake, any given scientific hypothesis. There is no general agreement in science about how to do this, which motivates putting forward a logical formulation of the problem. We do so by investigating distinct definitions of ‘rejection degrees’ each yielding a consequence relation. Our investigation leads to novel variations on the theme of *rational consequence relations*, prominent among non-monotonic logics.

*Keywords:* data-driven inference; significance inference; null hypothesis significance testing; non-monotonic logic

## 1 Introduction and motivation

The research reported in this note originates in [2], where a case is made for investigating consequence relations capable of expressing certain aspects of scientific inference. In particular, we are interested in capturing the fact that data may lead scientists to reject, possibly to some degree, and possibly by mistake, any given scientific hypothesis.

The canonical treatment of this general problem makes one of its first appearances in a 1925 book, which went on to shape the methodology of much empirical science [17]. In it, R.A. Fisher draws attention to ‘the logical nature’ of the inference underpinning *tests of significance*. The problem he sets out to address is that of examining a scientific hypothesis ( $H_0$ ) based on the observations it leads to, if true. When the observations so obtained are improbable enough, they provide grounds for us to reject the assumption that  $H_0$  is indeed true. Fisher’s line of reasoning culminates with what became known as the *Fisher disjunction*:

\*E-mail: paolo.baldi@unisalento.it

\*\*E-mail: esther.corsi@unimi.it

†E-mail: hykel.hosni@unimi.it

## 2 A logical framework for data-driven reasoning

The force with which such a conclusion is supported is logically that of the simple disjunction: *Either* an exceptionally rare event has occurred, *or*  $[H_0]$  is not true. ([18], p.39.)

We understand Fisher's 'simple' as 'being governed by logic', which of course for him could only be (an informal version of) classical logic. Setting aside some rare yet notable exceptions discussed in Section 5.1, the statistical and methodological communities seem to have taken this informal-classical-logic view at face value. Interestingly, this applies to both supporters and critics of the procedures put forward by Fisher. For the focus of the long-standing debates on tests of significance, and in particular on the infamous p-value [58] is usually on the meaning and nature of probability, rather than on the properties of the inferences scientists make with it.

This paper takes a logical perspective on inference based on the data-driven rejection of scientific hypotheses. And, insofar as this is possible, it aims to do so independently of any philosophical view on probability. We ask which properties are desirable for a consequence relation whose intended semantics is based on the degree of rejection that a given set of data provides to a given hypothesis. Hence we consider a more abstract and general problem than that envisaged by Fisher (and his many followers), which is however recovered as a special case of our question. Our results suggest that reasonable answers to our main question will be variations on the theme of non-monotonic consequence relation, in the sense brought to the attention of logicians by [30, 34, 36, 52].

The remainder of this introductory section provides the essential background on the kind of scientific inference we focus on, along with the required logical preliminaries. We clearly do not aim at exhausting the many ramifications of the vast topic of scientific inference. On the contrary, the following two subsections should be thought of as delimiting the scope of our investigation. We defer the discussion on related work to the concluding section of the paper.

### 1.1 NHST

According to a basic tenet in the methodology of science, to assess a scientific hypothesis, one looks at its logical consequences. In this context, the pattern of inference captured by *modus tollens* becomes prominent. If a sentence  $\theta$ , which is known to be false, follows logically from sentence  $\varphi$ , then we should conclude  $\neg\varphi$ . As we will now recall, this classical pattern of inference is often taken to lend its validity to the procedure known as Null Hypothesis Statistical Testing (NHST).

Let  $H_0$  be a sentence (in classical logic) standing for a scientific hypothesis, usually referred to as the *null hypothesis*. The key idea in NHST is to set up an experiment that leads to observations under the assumption that  $H_0$  is true. Denote by  $\sigma$  a *test statistics*, i.e. a function of the observations thus obtained. Finally, associate to this function the quantity usually referred to as the 'observed level of significance', or the 'probability value', or simply the 'p-value'. This latter is the calculated conditional probability of seeing equally extreme or more extreme data (according to the test statistics) given  $H_0$ . In analogy with *modus tollens* one says that  $H_0$  should be rejected if the p-value is small enough. Here is a schematization of the argument:

- (1) Suppose  $H_0$ ;
- (2) Calculate the p-value for some test statistics  $\sigma$  (i.e. a well-defined function of the data observed conditional on  $H_0$ );
- (3) The smaller the p-value, the stronger the reason to believe that  $H_0$  is not true.

Many statisticians and methodologists explicitly draw the parallel between (versions of) the above and *modus tollens*. To make a few notable examples, the authors of [13] speak of 'inductive *modus tollens*', [39] refers to 'statistical *modus tollens*', whereas [51] tellingly notes that it is the analogy with *modus tollens* that gives tests of significance prominence in the 'scientific method':

[Modus tollens] is at the heart of the philosophy of science, according to Popper. Its statistical manifestation is in [the] formulation of hypothesis testing that we will call ‘rejection trials’. ([51], p.72)

Quite interestingly, this view is shared also by prominent critics of NHST, e.g. [55].

While the appeal of the loose analogy with modus tollens is clear, it is misleading nonetheless, for no probabilistic test will deliver  $\neg H_0$ . Royall is again an authoritative voice who acknowledges the problem, but then argues that it is of no real consequence. ‘But the form of reasoning in the statistical version of the problem parallels that in deductive logic: if  $H_0$  implies  $E$  (with high probability), then not- $E$  justifies rejecting  $H_0$ ’ [51], p.73.

The qualitative difference between  $H_0$  being classically false and it being very improbable did not escape the attention of Fisher and his scrupulous followers, when they insist that the p-value is best interpreted as the degree to which observational data turns out to be incompatible with  $H_0$ . So a significance test can only lead one to conclude that  $H_0$  is *unlikely to be true*, if the p-value is small enough. But then the validity of this conclusion owes to the Fisher disjunction, rather than to modus tollens. And in turn, the former is grounded on the metaphysical assumption that small-probability events normally do not happen. Many methodologists, probabilists and statisticians strongly disagree with this, as testified by the animosity of the long-lasting debate on this topic [26, 58, 59].

Since ‘being unlikely true’ is logically distinct from ‘being classically false’, we have no reason to believe that the above schematization holds under uncertainty. Indeed, it is well known that modus tollens need not be adequate for probabilistic reasoning [6], and in general fails to deliver a point-valued probability [57]. Given that uncertainty is the norm, rather than the exception in science, and that uncertainty is most typically quantified probabilistically, the standard justification for adopting the canonical significance tests is, at the very least, in clear need of a logical clarification.

Note that a line of criticism has indeed addressed ‘the logic of NHST’, concluding that it is defective [15, 31, 46, 55]. The following example, which probably appears for the first time in [46], is by now standard.

#### EXAMPLE 1.1

Consider the following argument.

- 1 Either Harold is a US citizen or he is not;
- 2 If Harold is a US citizen, then he is most probably not a member of Congress;
- 3 Harold is a member of Congress;
- ∴ Harold is likely not a US citizen.

While assumptions 1–3 are all true, the conclusion is absurd, since being a US citizen is a necessary condition to sit in the US Congress. And yet it would be arrived at through (a kind of) modus tollens.

Those who use it take the argument in Example 1.1 to be conclusive in showing the fallacious nature of NHST. Note however that the criticism is conclusive only if one assumes that classical logic is the logic against which a pattern of inference can be evaluated. However, as remarked above, classical logic is hardly adequate to capture the ‘most probably’ and the ‘likely’, which appear in 2 and in the conclusion. Non-monotonic consequence relations on the contrary can express them, as detailed in Section 4.3. And indeed, if one looks at it from the non-monotonic-logic point of view, a natural conclusion of premisses 1–3 in Example 1.1 is that *Harold is not a typical US citizen*.

Finally a remark on terminology. Statisticians and methodologists go to great lengths to distinguish significance tests from the associated decision of either ‘retaining’ or ‘rejecting’ the null hypothesis,

## 4 A logical framework for data-driven reasoning

a procedure made prominent by J. Neyman and E. Pearson. In this latter approach one fixes, before running the relevant experiment, a given threshold for rejection, usually referred to as ‘critical region’ and denoted by  $\alpha$ . If, conditional on  $H_0$ , the observations fall within this region, then  $H_0$  is rejected. However, the noticeable mathematical commonalities between the Fisher and the Neyman–Pearson takes on the problem lead to the mishmash of the two methodologies under the heading NHST, see [12] for a comprehensive analysis. In what follows we can be rather casual about this internal divide within the classical statistical tradition, because none of our results below is affected by committing to either of the competing views.

### 1.2 Strong inference

The method of *strong inference* was put forward in 1964 *Science* editorial by mathematician and biophysicist John Platt [47]. Ever since it has been argued, especially in the life sciences, to provide a methodological guidance to hypothesis testing for working scientists [28] and indeed a sound generalization of NHST [4].

Platt offers the following schematic representation of strong inference in [47]:

- 1) Devising alternative hypotheses;
- 2) Devising a crucial experiment (or several of them), with alternative possible outcomes, each of which will, as nearly as possible, exclude one or more of the hypotheses;
- 3) Carrying out the experiment so as to get a clean result;
- 4) Recycling the procedure, making subhypotheses or sequential hypotheses to refine the possibilities that remain; and so on.

According to the author, the distinctive trait of strong inference lies in the fact that it requires scientists to produce as many competing hypotheses as possible in the given context. This gives rise to a rich tree of alternatives. As strong inference proceeds, the tree is pruned by (the logical consequences of) observations, i.e. outcomes of ‘crucial experiments’. The process runs until, in the best possible scenario, a single hypothesis survives.

Strong inference appears to have attracted very little interest from the logical community. However, we suggested in [2] that, owing to the conceptual analogies with the Ulam–Rényi game, introduced by Alfréd Rényi [48] and Stanislaw Ulam [56], strong inference can give us useful cues about the desirable properties of a consequence relation formalizing data-driven inference.

### 1.3 Non-monotonic consequence relations

We will be using several symbols for related but distinct consequence relations. As usual,  $\models$  denotes classical (propositional) consequence, and  $\equiv$  classical equivalence. Then we use  $\vdash$  to denote an arbitrary non-classical consequence relation.

Given the revisable nature of scientific inference, our logical framework is closely related to the family of non-monotonic logics. Those originated in the 1980s, mainly within the artificial intelligence community, for knowledge representation and reasoning purposes, see [38] for a recent and broad appraisal with a comprehensive bibliography. After a couple of decades since the first proposals, the community of non-monotonic logicians converged on three important logical systems: *cumulative* System C, the *preferential* System P and the *rational* System R. Table 1 presents the latter.

System C and System P are two subsets of System R. System C is defined by (REF), (LLE), (RWE) and (CMO), whereas System P also satisfies (AND) and (OR). In the interest of readability, we do not introduce distinct symbols for the three consequence relations characterizing those systems.

TABLE 1. System R

|                                                                                                          |                                                                                             |
|----------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| $\frac{}{\varphi \vdash \varphi}$ (REF)                                                                  | $\frac{\varphi \vdash \psi \quad \psi \models \xi}{\varphi \vdash \xi}$ (RWE)               |
| $\frac{\varphi \vdash \xi \quad \varphi \models \psi \quad \psi \models \varphi}{\psi \vdash \xi}$ (LLE) | $\frac{\varphi \vdash \psi \quad \varphi \vdash \xi}{\varphi \wedge \xi \vdash \psi}$ (CMO) |
| $\frac{\varphi \vdash \psi \quad \varphi \vdash \xi}{\varphi \vdash \psi \wedge \xi}$ (AND)              | $\frac{\varphi \vdash \xi \quad \psi \vdash \xi}{\varphi \vee \psi \vdash \xi}$ (OR)        |
| $\frac{\varphi \vdash \psi \quad \varphi \not\vdash \neg \psi}{\varphi \wedge \xi \vdash \psi}$ (RMO)    |                                                                                             |

#### 1.4 Plan of the paper

Section 2 sets out the intended semantics for consequence relations based on ‘degrees of rejection’. Its formalization, denoted by  $\vdash_{RJ}$ , will constitute a blueprint for defining consequence relations that capture specific aspects of this. Section 3 investigates  $\vdash_{IRJ}$  and  $\vdash_{\alpha RJ}$ , which are based on probabilistic rejection and are closely related to maximum likelihood and NHST, respectively. Then, building closer ties with USI games, we introduce and investigate  $\vdash_{uRJ}$  in Section 4. This satisfies (UMO), a form of constrained monotonicity, which, to the best of our knowledge, has not been investigated before. Theorem 4.14, the main result of this paper, establishes a form of completeness for  $\vdash_{uRJ}$ . Section 5 summarizes our contributions and lists future works on this topic.

## 2 The blueprint RJ-consequence

The key contribution of this paper is a proposal for the formalization of consequence relations capturing a central feature of scientific inference: rejecting certain hypotheses in the light of data. Doing this requires inevitable abstraction so, to pin down the intended semantics of the consequence relations of interest, we first strip down our problem of many of the details encountered in scientific practice. This, we submit, leaves us with three central features of data-driven inference: (i) there is a distinction between data and hypotheses, but it pertains to the attitude scientists have towards the statements representing them, rather than the nature of the statements themselves; (ii) experiments (the data-generating processes) can be repeated, and may lead to non-unique outcomes; (iii) when making data-driven inferences, scientists can rely on knowledge that is not subject to questioning by the experiment at hand. Since Definition 2.1 below captures the interaction of those three features, let us first motivate their desirability in the light of data-driven inference. While doing so, we will also introduce some terminology and notation.

To tackle the subtleties involved in (i) we work with a language composed of two (finite) sets of propositional variables  $\mathcal{H} = \{H_1, \dots, H_n\}$  and  $\mathcal{D} = \{d_1, \dots, d_l\}$ , standing for *hypotheses* and *data*, respectively. We do not assume that the sets are disjoint, and sometimes we require them not to be, as in Proposition 2.7 below. We denote by  $\text{Fm}_{\mathcal{H}}$  and  $\text{Fm}_{\mathcal{D}}$  the set of propositional formulas generated by closing  $\mathcal{H}$  and  $\mathcal{D}$ , respectively, under the usual connectives  $\wedge, \vee, \neg$ . The set of all formulas is denoted by  $\text{Fm}$ , i.e.  $\text{Fm} = \text{Fm}_{\mathcal{D}} \cup \text{Fm}_{\mathcal{H}}$ . We denote elements of  $\text{Fm}$  by lowercase Greek letters.

## 6 A logical framework for data-driven reasoning

The key element in our construction is *the degree to which data  $\delta \in \text{Fm}_{\mathcal{D}}$  rejects a hypothesis  $\varphi \in \text{Fm}_{\mathcal{H}}$* . As will be clear in Section 3 and in Section 4, detailing the specific properties of degrees of rejection will lead to distinct consequence relations. This is why we refer to the Definition 2.1 below as a *blueprint* consequence relation.

To motivate (ii) recall that one key feature in the complicated process that leads to data-driven scientific knowledge is the ‘replicability’ of the tests or experiments that engender it. The relevance of this for the intended semantics of the consequence relations we are about to construct is that the same hypothesis may be confronted multiple times by the same data. A textbook situation in which this happens is when sampling with replacement from an urn. But of course, concrete cases of scientific inference will not be so easy to describe (mathematically). Moreover, scientists may make mistakes in the experimental setup, in the data analysis, or at any other point of the procedure. We accommodate those heterogeneous cases centring Definition 2.1 on the degree of rejection that a *multiset* of data  $\Delta$  from  $\text{Fm}_{\mathcal{D}}$  puts forward against hypothesis  $\varphi \in \text{Fm}_{\mathcal{H}}$ .

We denote by  $\mathcal{M}_{\mathcal{D}}$  the set of multisets built over the formulas in  $\text{Fm}_{\mathcal{D}}$ . Henceforth, abusing notation we will use curly brackets for denoting both a multiset and a set of elements. Recall that the set  $\mathcal{M}_S$  of multisets built over a set  $S$  can be more formally identified with the set  $\mathbb{N}^S$  of functions from  $S$  to  $\mathbb{N}$ , where  $f \in \mathbb{N}^S$  stands for the multiset where each  $s \in S$  occurs  $f(s) = n$  times.

Finally, to motivate (iii), note that in scientific inference data is not *certain*. This is one reason why the replicability of experimental setups is of critical importance. However, since we are putting forward a logical framework, we must provide our ideal and logically omniscient scientists with knowledge about the problem at hand, which is *not* revisable as a consequence of the experiment. To capture this we will assume that data-driven inference is performed under the assumption that the hypotheses in  $\mathcal{H} = \{H_1, \dots, H_n\}$  are mutually exclusive and exhaustive. We capture this through the following classical formula  $T$

$$T := \bigvee_{H \in \mathcal{H}} H \wedge \bigwedge_{H, H' \in \mathcal{H}, H \neq H'} H \rightarrow \neg H'.$$

We are now ready to provide the blueprint definition of the degree  $r$  to which data reject a hypothesis. Sections 3 and 4 will investigate two special cases that arise by specifying distinct ways of computing  $r$ .

### DEFINITION 2.1

We say that  $r: \mathcal{M}_{\mathcal{D}} \times \mathcal{H} \rightarrow \mathbb{R}_+ \cup \{\infty\}$  is a *degree of rejection* if, for any  $\gamma, \delta \in \text{Fm}_{\mathcal{D}}$ ,  $\Delta \in \mathcal{M}_{\mathcal{D}}$  and  $H \in \mathcal{H}$ , the following hold:

- (1) If  $T, \gamma \models \delta$ , then  $r_{\gamma}(H) \geq r_{\delta}(H)$ .
- (2) If  $T, H \models \delta$  then  $r_{\delta}(H) = 0$ .
- (3) If  $T, \delta \models \neg H$ , then  $r_{\delta}(H) > 0$ .
- (4) If  $r_{\Delta}(H) \neq \infty$ , then  $r_{\Delta}(H) = \sum_{\delta \in \Delta} r_{\delta}(H)$ .

When no confusion is likely to arise, we simply speak of  $r_{\Delta}$  as ‘the degree of rejection’ and denote it by  $r$ .

Proposition 3.1 below shows that the definition is well-posed. In fact the remainder of this paper provides examples of functions that satisfy Definition 2.1 and explores how they provide the semantics for distinct consequence relations.

Part (1) of the definition states that, modulo  $T$ , logically stronger data provide a higher degree of rejection. If we do not insist on limiting our inferences to satisfiable data, this implies that observing contradictions provides the maximal degree of rejection for any hypothesis in  $\text{Fm}_{\mathcal{H}}$ . This is at odds

with the scientific practice where we may observe a contradiction between the data and a hypothesis, but we typically do not observe contradictory data. As a consequence, we will tacitly assume that the elements of  $\text{Fm}_{\mathcal{D}}$  are satisfiable even if the formalism does not require us to do so.

Conditions (2) and (3) link degrees of rejection to classical logic. The former states that no positive degree of rejection can be provided when one observes a logical consequence, modulo  $T$ , of a given hypothesis. (3) states that a degree of rejection cannot be zero whenever the data, modulo  $T$ , entail the negation of the hypothesis. This is where Definition 2.1 marks a genuine departure from the traditional idea that the relation between observations and hypotheses is governed entirely by classical deduction, and in particular through *modus tollens*. While it may happen, in special cases, that  $r_{\delta}(H) = 1$  for data and hypotheses such that  $\delta \models \neg H$ , this need not be true in general and  $r$  can take any positive value up to  $\infty$ . This departure from classical logic is motivated, as noted in the introductory section of this work, by the fact that data obtained in any experimental setup may translate into incomplete, scarce, noisy or otherwise imperfect evidence. Condition (3) ensures that the semantics of data-driven inference accommodates this central feature of scientific inference at root.

Finally, condition (4) states that the degree to which a hypothesis is rejected by data is computed as the sum of the degrees of rejection obtained through each single piece of data in the multiset  $\Delta$ . As a justification for aggregating with the sum, as opposed to any other function, at present we can offer its simplicity and the interesting consequences it leads to. Note however that additive aggregations are used across many (distinct) formalisms in uncertain reasoning, including the sums of ‘uncertainties’ in Adams’s probabilistic logic [1], and the sums of ‘risks’ in Giles’s game-theoretic foundation of many-valued logic [21].

The next step towards the intended semantics of our blueprint consequence relation revolves around the set  $\hat{\mathcal{H}}_{\Delta}$  of hypotheses in  $\mathcal{H}$ , which are *least rejected* by the data in  $\Delta$ :

$$\hat{\mathcal{H}}_{\Delta} = \arg \min_{H \in \mathcal{H}_{\Delta}^f} r_{\Delta}(H), \quad (1)$$

where  $\mathcal{H}_{\Delta}^f$  is the set of hypotheses with a finite degree of rejection. The rationale behind least-rejection is that the elements of  $\hat{\mathcal{H}}_{\Delta}$  are those hypotheses that ‘did best against data’ in  $\Delta$ .

#### REMARK 2.2

$\hat{\mathcal{H}}_{\Delta}$  is empty only when there is no hypothesis with a finite rejection degree, i.e. if  $r_{\Delta}(H) = \infty$  for each  $H \in \mathcal{H}$ .

We are now ready to define the blueprint consequence  $\vdash_{RJ}$  as classical model-preservation under least-rejection.

#### DEFINITION 2.3 (RJ-consequence).

Let  $\Delta$  be a multiset of formulas in  $\text{Fm}_{\mathcal{D}}$  and  $\varphi \in \text{Fm}_{\mathcal{H}}$ . We say that  $\varphi$  is an *RJ-consequence* of  $\Delta$ , written  $\Delta \vdash_{RJ} \varphi$ , if  $T, H \models \varphi$ , for each  $H \in \hat{\mathcal{H}}_{\Delta}$ .

The instantiations of the blueprint consequence to be investigated in this paper will make explicit the nature of the ordering among rejected hypothesis that makes the minimality featuring in (1) precisely defined. But readers who are familiar with the preferential variety of non-monotonic logic will realize immediately from Definition 2.1 that our blueprint belongs to that family.

As an immediate consequence of Definition 2.3 and Remark 2.2 we have that  $\vdash_{RJ}$  is explosive in the sense that if there are no least-rejected hypotheses, then all hypotheses are rejected by data.

8 *A logical framework for data-driven reasoning*TABLE 2. System  $RJ$ 

$$\frac{\Delta \vdash \varphi \quad \varphi \models \psi}{\Delta \vdash \psi} \text{ (RWE)}$$

$$\frac{\Delta, \gamma \vdash \varphi \quad \models \gamma \leftrightarrow \delta}{\Delta, \delta \vdash \varphi} \text{ (LLE)} \quad \frac{\Delta \vdash \varphi \quad \Delta \vdash \psi}{\Delta \vdash \varphi \wedge \psi} \text{ (AND)}$$

## LEMMA 2.4

Suppose  $\hat{\mathcal{H}}_\Delta = \emptyset$ . Then  $\Delta \vdash_{RJ} \varphi$  for any  $\varphi \in \text{Fm}_{\mathcal{H}}$ .

Proposition 2.5 shows that  $\vdash_{RJ}$  satisfies (multiset versions of) the rules (RWE), (LLE) and (AND), whereas Proposition 2.6 shows that  $\vdash_{RJ}$  fails, desirably, unconstrained monotonicity. Finally, Proposition 2.7 shows that  $\vdash_{RJ}$  satisfies the rules of rational consequence relation with the important exception of (OR), which will be discussed later.

## PROPOSITION 2.5

$RJ$ -consequence satisfies the rules in Table 2.

See the [Appendix](#) for the proof.

It is well-known [23, 36, 37] that (non-monotonic) logical systems satisfying (AND) provide *qualitative* representations of reasoning under uncertainty. This may appear to be at odds with the ‘numerical’ semantics provided by degrees of rejection. However this is not sufficient to make  $\vdash_{RJ}$  quantitative. In fact the magnitude of degrees of rejection are immaterial in (1), where only comparisons of such degrees are relevant.

Next, we ensure that  $\vdash_{RJ}$  does not satisfy unconstrained monotonicity (MON). Since our consequence relations are defined on multisets of premisses, two distinct formulations of unconstrained monotonicity must be taken into account:

$$\frac{\gamma \vdash \varphi}{\gamma \wedge \delta \vdash \varphi} \text{ (AMON)} \quad \frac{\Delta \vdash \varphi}{\Delta, \Delta' \vdash \varphi} \text{ (MMON)},$$

where, as usual,  $\vdash$  stands for an arbitrary consequence relation.

## PROPOSITION 2.6

Neither (AMON) nor (MMON) hold for  $\vdash_{RJ}$ .

See the [Appendix](#) for the proof.

In preparation for the last proposition of this Section, which pins down the conditions under which reflexivity, rational monotonicity and cut are satisfied by  $\vdash_{RJ}$ , recall that we work with distinct languages for data and hypotheses. Since Proposition 2.5 involves only rules that do not feature individual formulas on both sides of  $\vdash_{RJ}$ , the proof carries through even if the two languages are disjoint. This possibility, however, must be ruled out for  $\vdash_{RJ}$  to satisfy the rules of the next Proposition. Moreover we must set to infinity the rejection degree provided by contradictory data against any hypothesis. Recall also that we denote by  $T$  the formula expressing that hypotheses form a partition.

TABLE 3.

|                                                                                            |                                                                                                     |
|--------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------|
| $\frac{}{\varphi \sim \varphi}$ (REF)                                                      | $\frac{\Delta \vdash \psi \quad \Delta \not\vdash \neg \varphi}{\Delta, \varphi \vdash \psi}$ (RMO) |
| $\frac{\Delta, \varphi \vdash \psi \quad \Delta \vdash \varphi}{\Delta \vdash \psi}$ (CUT) | $\frac{\Delta \vdash \psi \quad \Delta \vdash \varphi}{\Delta, \varphi \vdash \psi}$ (CMO)          |

**PROPOSITION 2.7**

Suppose  $\text{Fm}_{\mathcal{D}} \cap \text{Fm}_{\mathcal{H}} \neq \emptyset$ . Assume further that if  $T \models \neg \delta$  then  $r_{\delta}(H) = \infty$  for each  $H \in \mathcal{H}$ . Then  $\sim_{RJ}$  satisfies the rules in Table 3.

See the [Appendix](#) for the proof.

For  $\sim_{RJ}$  to be a rational consequence relation it should also satisfy (OR). It does not, as detailed in Section 3.2, where we also explain why this is indeed desirable.

As a concluding remark on our blueprint consequence relation, note that the *supraclassicality* of  $\sim_{RJ}$  is an immediate corollary of Proposition 2.7 by concatenating the arguments for (REF) and (RWE).

**COROLLARY 2.8**

Let  $\varphi, \psi \in \text{Fm}_{\mathcal{D}} \cap \text{Fm}_{\mathcal{H}}$ . If  $\varphi \models \psi$  then  $\varphi \sim_{RJ} \psi$ .

### 3 Data-driven probabilistic rejection

This section investigates two ways of defining degrees of rejection through data-driven probabilities. The first is discussed in Section 3.1 and captures the core idea of maximum likelihood inference. The second is discussed in Section 3.3 and is closely related to significance inference and NHST. Each gives rise to a distinct consequence relation. We argue that both capture interesting aspects of statistical inference.

Recall that a *probability function* on  $\text{Fm}$  is a map  $P : \text{Fm} \rightarrow [0, 1]$ , which satisfies normalization and additivity, i.e.

- (1) If  $\models \varphi$  then  $P(\varphi) = 1$  and
- (2) If  $\models \neg(\varphi \wedge \theta)$  then  $P(\varphi \vee \theta) = P(\varphi) + P(\theta)$ .

Throughout this section we will work with *conditional* probability functions of  $\text{Fm}$ , i.e.  $P : \text{Fm} \times \text{Fm}^+ \rightarrow [0, 1]$  defined by

$$P(\delta \mid H) = \frac{P(\delta \wedge H)}{P(H)}, \quad (2)$$

where  $\text{Fm}^+$  denotes the set of elements of  $\text{Fm}$  whose probability is strictly positive. Formally, this restriction guarantees we do not divide by 0. Conceptually, it prevents bizarre applications of statistical hypothesis testing where the hypotheses being tested are known to be false.

## 10 A logical framework for data-driven reasoning

Functions defined as in (2) above are normalized on tautologies and finitely additive. That is, they satisfy:

$$\text{If } \models \neg(\delta \wedge \gamma) \text{ then } P(\neg\delta \vee \neg\gamma \mid H) = P(\delta \mid H) + P(\gamma \mid H). \quad (3)$$

To express the condition to the effect that  $\delta$  and  $\gamma$  are incompatible, we will also write  $\delta, \gamma \models \perp$ , where  $\perp$  is any classical contradiction. Finally the following holds and will be used extensively below:

$$P(\neg\delta \mid H) = 1 - P(\delta \mid H). \quad (4)$$

### 3.1 IRJ-consequence

To illustrate the idea of likelihood-based rejection, consider  $\mathcal{H} = \{H_1, H_2\}$  and let

$$r_\delta^l(H_i) = P(\neg\delta \mid H_i) \quad (5)$$

for data  $\delta \in \text{Fm}_{\mathcal{D}}$  and  $i \in \{1, 2\}$ . In this situation  $H_1$  belongs to the set of least rejected hypotheses  $\hat{\mathcal{H}}_\delta$  if and only if  $r_\delta(H_1) \leq r_\delta(H_2)$ . By (4) then  $H_1 \in \hat{\mathcal{H}}_\delta$  if and only if  $1 - P(\delta \mid H_1) \leq 1 - P(\delta \mid H_2)$ , i.e.

$$H_1 \in \hat{\mathcal{H}}_\delta \text{ if and only if } P(\delta \mid H_1) \geq P(\delta \mid H_2).$$

So, when  $r$  is computed as in (5), the set of least rejected hypotheses under  $\delta$  pins down the set of hypotheses that maximize likelihood. Hence we call  $r^l$  the *maximal likelihood rejection function*.

As above we assume the additive aggregation of the degrees of rejection over multisets of data:

$$r_\Delta^l(H) = \sum_{\delta \in \Delta} r_\delta^l(H).$$

#### PROPOSITION 3.1

$r_\delta^l(H)$  is a degree of rejection.

PROOF. We must show that  $r^l$  satisfies parts (1-3) of Definition 2.1. For part (1), recall that  $\gamma \models \delta$  entails  $P(\gamma \mid H) \leq P(\delta \mid H)$  for all conditional probability functions defined as in (2). By (4), we get  $r_\delta^l(H) = P(\neg\delta \mid H) \leq P(\neg\gamma \mid H) = r_\gamma^l(H)$ , as required. As for (2) observe that if  $H \models \delta$  then  $P(\delta \mid H) = 1$ . By (4) we have  $r_\delta^l(H) = P(\neg\delta \mid H) = 1 - P(\delta \mid H) = 0$ , as required. A similar argument delivers (3). For if  $H \models \neg\delta$ , we have  $r_\delta^l(H) = P(\neg\delta \mid H) = 1 - P(\delta \mid H) = 1$ , which is greater than 0.  $\square$

We are now ready to define  $\vdash_{IRJ}$ , which arises from  $\vdash_{RJ}$  by imposing  $r^l$  as the degree of rejection. Recall that  $\Delta \vdash_{RJ} \varphi$  if  $\varphi$  is a classical consequence of  $H$  for each minimally rejected hypothesis  $H \in \hat{\mathcal{H}}_\Delta$ .

#### DEFINITION 3.2

We say that  $\varphi$  is a *IRJ-consequence* of  $\Delta$ , written  $\Delta \vdash_{IRJ} \varphi$ , if

- i)  $\Delta \vdash_{RJ} \varphi$ , and
- ii)  $r$  is computed according to (5), i.e.  $r = r^l$ .

As a consequence of Proposition 3.1,  $\vdash_{IRJ} \supset \vdash_{RJ}$ . The key difference between  $\vdash_{IRJ}$  and the blueprint  $\vdash_{RJ}$  lies in the fact that the former satisfies both *conjunction on the left* and its *converse*,

namely:

$$\frac{\Delta, \delta, \gamma \vdash \varphi \quad \neg \gamma, \neg \delta \models \perp}{\Delta, \delta \wedge \gamma \vdash \varphi} (\text{AND})_I$$

and

$$\frac{\Delta, \delta \wedge \gamma \vdash \varphi \quad \neg \gamma, \neg \delta \models \perp}{\Delta, \delta, \gamma \vdash \varphi} (\text{AND})_I^{\text{con}}$$

while this is not the case in general for  $\vdash_{RJ}$ , as we will show in the next section.

Before proving that  $(\text{AND})_I$  holds for  $\vdash_{IRJ}$ , let us argue in favour of its desirability. By inspecting the likelihood-based rejection function  $r^I$  (5) it is apparent that the semantics of  $\vdash_{IRJ}$  builds on ‘the probability of negations’. It is therefore the additivity of conditional probability functions (3) that delivers the validity of this rule.

### PROPOSITION 3.3

In addition to all the rules in Table 2  $IRJ$ -consequence relations satisfy also  $(\text{AND})_I$  and  $(\text{AND})_I^{\text{con}}$ .

PROOF. Suppose  $\Delta, \delta, \gamma \vdash \varphi$  and  $\neg \delta, \neg \gamma \models \perp$ . By (5), we have  $r_{\delta \wedge \gamma}^I(H) = P(\neg(\delta \wedge \gamma) \mid H)$ . The right hand side, by the elementary properties of conditional probability functions equals  $P(\neg \delta \vee \neg \gamma \mid H)$ . By the second hypothesis of  $(\text{AND})_I$  we know that  $\neg \delta$  and  $\neg \gamma$  are incompatible. Hence, by additivity  $r_{\delta \wedge \gamma}^I(H) = P(\neg \delta \mid H) + P(\neg \gamma \mid H)$ , which equals  $r_{\delta}^I(H) + r_{\gamma}^I(H)$ . This entails that for each  $H$ , we have  $r_{\Delta, \gamma, \delta}^I(H) = r_{\Delta, \gamma \wedge \delta}^I(H)$ , hence  $\hat{\mathcal{H}}_{\Delta, \gamma, \delta} = \hat{\mathcal{H}}_{\Delta, \gamma \wedge \delta}$  and thus  $\Delta, \gamma, \delta \vdash_{RJ} \varphi$  if and only if  $\Delta, \gamma \wedge \delta \vdash_{RJ} \varphi$ . This finally shows that  $r^I$  satisfies both  $(\text{AND})_I$  and  $(\text{AND})_I^{\text{con}}$ .  $\square$

### 3.2 (OR) is not valid for $\vdash_{IRJ}$

As a consequence of Proposition 3.1, if a rule is not valid for  $\vdash_{IRJ}$ , then it is not valid for  $\vdash_{RJ}$  either. And, as anticipated, a noticeable failure for the blueprint consequence relation is the rule known in the non-monotonic logic literature as *disjunction in the premisses* (OR). Hence both  $\vdash_{IRJ}$  and its predecessor  $\vdash_{RJ}$ , fall short of being preferential consequence relations in the sense recalled in Section 1.3.

Before showing this, let us make explicit that this *is* in fact desirable in light of the intended semantics of the consequence relation(s). Recall that  $\vdash_{IRJ}$  is likelihood based: the degree to which an observation  $\delta$  rejects a hypothesis  $H$  is computed as  $P(\neg \delta \mid H)$ . If we were to insist on satisfying (OR) we would need to guarantee that  $P(\neg(\delta_i \vee \delta_j) \mid H) = P(\neg \delta_i \wedge \neg \delta_j \mid H)$  can always be computed in terms of  $P(\neg \delta_i \mid H)$  and  $P(\neg \delta_j \mid H)$ . But it is well known that is cannot be done in general, for (conditional) probability functions are not compositional with respect to conjunction—see Section 1.1 of [25] for a general discussion.

We show that (OR) fails for  $IRJ$ -consequence by arguing that the weaker rule (XOR), introduced in [23], fails too:

$$\frac{\Sigma, \delta \vdash \varphi \quad \Sigma, \gamma \vdash \varphi \quad \delta, \gamma \models \perp}{\Sigma, \delta \vee \gamma \vdash \varphi} (\text{XOR})$$

### LEMMA 3.4

(XOR) does not hold for  $\vdash_{IRJ}$ .

## 12 A logical framework for data-driven reasoning

PROOF. To construct the required counterexample, let  $\mathcal{H} = \{H_1, H_2, H_3\}$  and suppose  $\gamma, \delta \models \perp$ , i.e.  $\gamma$  and  $\delta$  are logically incompatible data. Consider the following probability assignments:

$$\begin{aligned} P(\neg\delta \wedge \gamma \mid H_1) &= 0.5 & P(\neg\delta \wedge \gamma \mid H_2) &= 0.3 & P(\neg\delta \wedge \gamma \mid H_3) &= 0.7 \\ P(\delta \wedge \neg\gamma \mid H_1) &= 0.4 & P(\delta \wedge \neg\gamma \mid H_2) &= 0.5 & P(\delta \wedge \neg\gamma \mid H_3) &= 0.1 \\ P(\neg\delta \wedge \neg\gamma \mid H_1) &= 0.1 & P(\neg\delta \wedge \neg\gamma \mid H_2) &= 0.2 & P(\neg\delta \wedge \neg\gamma \mid H_3) &= 0.2. \end{aligned}$$

Since  $\gamma, \delta \models \perp$  we have

$$P(\delta \wedge \gamma \mid H_1) = P(\delta \wedge \gamma \mid H_2) = P(\delta \wedge \gamma \mid H_3) = 0.$$

Now, by simple computations, we obtain

$$\begin{aligned} r_\delta^l(H_1) &= P(\neg\delta \mid H_1) = 0.6, \\ r_\delta^l(H_2) &= P(\neg\delta \mid H_2) = 0.5, \\ r_\delta^l(H_3) &= P(\neg\delta \mid H_3) = 0.9, \end{aligned}$$

hence  $\hat{\mathcal{H}}_\delta = \{H_2\}$  and

$$\delta \sim H_2 \vee H_3.$$

Moreover,

$$\begin{aligned} r_\gamma^l(H_1) &= P(\neg\gamma \mid H_1) = 0.5, \\ r_\gamma^l(H_2) &= P(\neg\gamma \mid H_2) = 0.7, \\ r_\gamma^l(H_3) &= P(\neg\gamma \mid H_3) = 0.3, \end{aligned}$$

hence  $\hat{\mathcal{H}}_\gamma = \{H_3\}$  and

$$\gamma \sim H_2 \vee H_3.$$

We have thus shown that the premises of (XOR) are satisfied. On the other hand  $r_{\delta \vee \gamma}^l(H_1) = P(\neg\delta \wedge \neg\gamma \mid H_1) = 0.1$ ,  $r_{\delta \vee \gamma}^l(H_2) = P(\neg\delta \wedge \neg\gamma \mid H_2) = 0.2$  and  $r_{\delta \vee \gamma}^l(H_3) = P(\neg\delta \wedge \neg\gamma \mid H_3) = 0.2$ , so that  $\hat{\mathcal{H}}_{\delta \vee \gamma} = \{H_1\}$ , and finally

$$\delta \vee \gamma \not\sim H_2 \vee H_3$$

i.e. the conclusion of the rule (XOR) is not satisfied.  $\square$

### 3.3 $\alpha$ -RJ consequence

Recall that likelihood is the probability of *observed* data conditional on a (statistical null) hypothesis. It is a calculated, rather than estimated, quantity. A long-standing controversy in the methodology of statistics revolves around whether likelihood is all that one should take into account when making data-driven inferences, see [51] for an articulated account.

As recalled in the introductory Section, the practice known as Null Hypothesis Significance Testing (NHST) combines several aspects of statistical methodology and gives a prominent role to the p-value, which is the calculated probability of obtaining *hypothetical* data, namely data at least as improbable as those observed conditional on the null hypothesis. Clearly the p-value is not entirely

data-driven, as it features data that in fact one has *not* seen—see [58] for an extensive discussion on this and related point.

For our present purposes it is enough that we capture the qualitative comparison of the conditional probabilities involved, and in particular the fact that when computing p-values, we consider the cumulative probability function defined over all events that are less probable than those observed conditional on  $H_0$ .

To capture this new feature, we add to the language of  $\vdash_{RJ}$  and  $\vdash_{IRJ}$  an operator  $t_H$ , which associates to each hypothesis  $H$ , the formula  $t_H(\delta)$ . The operator is characterized as follows:

$$\delta \models t_H(\delta), \quad \frac{\gamma \models \delta}{t_H(\gamma) \models t_H(\delta)}, \quad t_H(t_H(\delta)) \models t_H(\delta). \quad (6)$$

From the logical point of view,  $t_H$  is just a closure operator over the set of models of the formula  $\delta \in \text{Fm}_{\mathcal{D}}$ . The rules that characterize it express its extensiveness, monotonicity and idempotence, respectively.

To illustrate,  $\delta \models t_H(\delta)$  states that the set of models of  $\delta$ , which translate probabilistically as the event that  $\delta$  occurs, is a subset of the set of models of  $t_H(\delta)$ , the event that data at least as improbable than  $\delta$  occur. The other rules read similarly.

The consequence relation to be introduced in Definition 3.7 below departs in another way from  $\vdash_{IRJ}$ . Recall that NHST combines ideas by Fisher about significance inference with ideas by Neyman and Pearson about statistical testing, albeit against the will of all parties involved. While the Fisher approach insists on p-values providing a degree of inconsistency between the data and the null hypothesis, the Neyman–Pearson one is oriented towards the (binary) decision whether to reject or not the null hypothesis. To do so, a threshold  $\alpha$  is fixed at design time, which identifies the so-called *rejection region*: the null hypothesis is rejected if the observations conditional on the null hypothesis falls within this region. The key feature of the consequence relation  $\vdash_{\alpha RJ}$  is that of combining the qualitative nature of the binary decision with the graded interpretation of the p-value.

Two preliminary remarks before we get to the definitions. First, we represent the p-value associated to hypothesis  $H$  by the quantity  $P(t_H(\delta) \mid H)$ . Building on this, the rejection degree for a hypothesis  $H \in \mathcal{H}$  is set to 1 minus the p-value, *provided* the latter is below a fixed threshold  $\alpha$ . Second, in Definition 3.5 below, and in the reminder of this Section, we will assume that hypotheses in  $\mathcal{H} = \{H_1, \dots, H_n\}$  are mutually exclusive and jointly exhaustive.

DEFINITION 3.5 ( $r^\alpha$ ).

Fix  $\alpha_i \in [0, 1)$  for each  $H_i \in \mathcal{H}$  and suppose  $\delta \in \text{Fm}_{\mathcal{D}}$ . Then

$$r_\delta^\alpha(H_i) = \begin{cases} 1 - P(t_{H_i}(\delta) \mid T \wedge H_i) & \text{if } P(t_{H_i}(\delta) \mid T \wedge H_i) \leq \alpha_i \\ 0 & \text{otherwise.} \end{cases}$$

Finally,

$$r_\Delta^\alpha(H) = \sum_{\delta \in \Delta} r_\delta^\alpha(H),$$

for any  $H \in \mathcal{H}$ ,  $\Delta \in \mathcal{M}_{\mathcal{D}}$ .

LEMMA 3.6

$r^\alpha$  is a degree of rejection.

PROOF. We need to check that the rejection degree of Definition 3.5 satisfies Definition 2.1.

## 14 A logical framework for data-driven reasoning

For condition (1), in suppose  $\gamma \models \delta$ , for  $\gamma, \delta \in \text{Fm}_{\mathcal{D}}$ . Since  $\gamma \models \delta$ , we have that  $t_{H_i}(\gamma) \models t_{H_i}(\delta)$ , hence by the monotonicity of probability functions,  $P(t_{H_i}(\gamma) \mid H_i) \leq P(t_{H_i}(\delta) \mid H_i)$ . From this it follows immediately that if  $P(t_{H_i}(\gamma) \mid H_i) \geq \alpha_i$ , then  $P(t_{H_i}(\delta) \mid H_i) \geq \alpha_i$ , and  $1 - P(t_{H_i}(\gamma) \mid H_i) \geq 1 - P(t_{H_i}(\delta) \mid H_i)$ , hence  $r_{\gamma}^{\alpha}(H) \geq r_{\delta}^{\alpha}(H)$ , as required.

The remaining conditions are established similarly, hence we omit checking them explicitly.  $\square$

We now define  $\sim_{\alpha RJ}$  as follows.

DEFINITION 3.7 ( $\sim_{\alpha RJ}$ ).

We say that  $\varphi$  is  $\alpha RJ$ -consequence of  $\Delta$ , written  $\Delta \sim_{\alpha RJ} \varphi$ , if

- i)  $\Delta \sim_{RJ} \varphi$ , and
- ii)  $r$  is computed according to Definition 3.5, i.e.  $r = r^{\alpha}$ .

The next example shows that  $\sim_{\alpha RJ}$  formalizes the key inferential step in a stylized example of a NHST procedure.

EXAMPLE 3.8

Consider the null hypothesis  $H_0$  and recall that we are under the blanket assumption to the effect that hypotheses form a partition. Hence we can write  $\mathcal{H} = \{H_0, \neg H_0\}$  instead of  $\mathcal{H} = \{H_0, H_1\}$ . We fix  $\alpha_0 \in (0, 1)$  and let  $\alpha_1 = 0$ . For simplicity, we drop the index of the former, denoting  $\alpha_0$  by  $\alpha$ . Definition 3.5 gives us  $r_{\delta}^{\alpha}(H_0) = 1 - P(t(\delta) \mid H_0)$  whenever  $P(t(\delta) \mid H_0) \leq \alpha$ , and  $r_{\delta}^{\alpha}(H_0) = 0$  otherwise. On the other hand  $r_{\delta}^{\alpha}(\neg H_0) = 1$  only if  $P(t_{\neg H_0}(\delta) \mid \neg H_0) = 0$ , and  $r_{\delta}^{\alpha}(\neg H_0) = 0$  otherwise. This accounts for the fact that in NHST one does not consider  $\neg H_0$  to be under scrutiny for possible rejection. Thus  $\neg H_0$  is rejected only in case of data practically impossible under  $\neg H_0$  (e.g. in case of logical contradictions), while the focus is wholly on whether the null hypothesis  $H_0$  is rejected. Hence, whenever  $P(t_{\neg H_0}(\delta) \mid \neg H_0) \neq 0$ , we have that  $\hat{\mathcal{H}}_{\delta} = \{\neg H_0\}$  if and only if  $r_{\delta}^{\alpha}(H_0) \neq 0$ , i.e.  $P(t(\delta) \mid H_0) \leq \alpha$ . This means in turn that

$$\delta \sim_{\alpha RJ} \neg H_0 \text{ if and only if } P(t(\delta) \mid H_0) \leq \alpha,$$

i.e. *the null hypothesis is rejected exactly when the p-value is below the fixed significance level (equivalently falls within the rejection region).*

Note that the above example also illustrates the asymmetry between *rejecting* the null hypothesis, i.e.  $\delta \sim_{\alpha RJ} \neg H_0$  and *retaining* it. Indeed, if  $P(t(\delta) \mid H_0) > \alpha$ , we have  $r_{\delta}^{\alpha}(H_0) = r_{\delta}^{\alpha}(\neg H_0) = 0$ , hence  $\hat{\mathcal{H}} = \{H_0, \neg H_0\}$ . So by Definition 3.7 we have both  $\delta \not\sim_{\alpha RJ} H_0$  and  $\delta \not\sim_{\alpha RJ} \neg H_0$ .

The last proposition in this Section compares  $\sim_{\alpha RJ}$  and  $\sim_{IRJ}$ , by showing that neither  $(\text{AND})_I$  nor  $(\text{AND})_I^{\text{con}}$ , which are satisfied by  $\sim_{\alpha RJ}$  (recall that by Proposition 3.3 they are both satisfied by  $\sim_{IRJ}$ ).

PROPOSITION 3.9

Both  $(\text{AND})_I$  and  $(\text{AND})_I^{\text{con}}$  are invalid under  $\sim_{\alpha RJ}$ .

See the Appendix for the proof.

## 4 uRJ-consequence

Our next and final instantiation of the blueprint consequence relation  $\sim_{RJ}$  arises by taking the more general approach to data-driven inference offered by the method of *Strong Inference*, briefly recalled in Section 1.2 above.

Strong Inference stands in close analogy with the Ulam–Rényi game, which we spell out as the *Ulam–Rényi game for Strong Inference*—USI, for short.

#### 4.1 Ulam–Rényi game for Strong Inference

A USI game is played by *Scientists* (**S**) against *Nature* (**N**). **N** ‘thinks’ of a number, called the *secret*, which is the index of the unique true hypothesis in  $\mathcal{H} = \{H_1, \dots, H_n\}$ . **S** must figure out what the secret is, which we denote by  $H_{s^*}$ , and aims to do so as quickly as possible. The only move available to **S** is to ask Nature binary questions, i.e. questions that can be answered with either ‘Yes’ or ‘No’. In this latter case we say that *the hypothesis has been rejected*.

Ulam–Rényi games are related to logic as follows. In the form just recalled above, they provide a sound and complete semantics for classical logic. If however, Nature is allowed to *lie*  $m$  times ( $m \in \mathbb{N}$ ), then USI game provides a sound and complete semantics for the  $(m + 2)$ -valued Łukasiewicz logic [41, 42].

Building on this, we will assume that **N** can lie *at most*  $0 \leq m < \infty$  times. We say that a hypothesis  $H$  is *temporarily rejected* if **N** answers ‘No’ to a question about  $H$  and  $m > 0$ . A hypothesis is *rejected* if it has been temporarily rejected at least  $m + 1$  times.

Rather than a deceptive or mischievous view of Nature, this feature of USI games captures the fact that data are often gappy, ambiguous or otherwise imperfect. Hence lies are interpreted here as mistakes made by scientists either in the design of experiments or in the analysis of the data generated by them. The fact that there is a bound to the number of mistakes captures the single most distinctive aspect of scientific reasoning: it will eventually self-correct. Indeed, the aim of the USI game is to reject all candidates except for the secret, which can be temporarily rejected, but not rejected.

The remainder of this work explores the consequences of taking the USI game as the intended semantics for our uRJ-consequence whose properties are pinned down in the main result of this paper, Theorem 4.14 below.

The formal setup is as follows. Each question-and-answer in a USI game is represented by a formula  $\varphi$  in Fm, which we interpret as a positive answer to the question ‘Does  $\varphi$  hold?’ We assume this is the only kind of data relevant to uRJ-consequence. So  $\Delta$  is now a multiset from  $\text{Fm} = \text{Fm}_{\mathcal{D}} = \text{Fm}_{\mathcal{H}}$ . This means that data is expressed in the same language of the hypotheses and we no longer distinguish  $\text{Fm}_{\mathcal{D}}$  from  $\text{Fm}_{\mathcal{H}}$ . As described in Section 2 scientists playing the USI game can rely on non-revisable assumptions, and in particular the fact that exactly one of the hypothesis is true (and Nature cannot lie about this). As usual we capture this with the formula  $T$ .

##### EXAMPLE 4.1

Suppose that  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$ . Recall that Scientists can only ask questions of the form ‘Does the secret belong to  $D \subseteq \mathcal{H}$ ?’ If, say,  $D = \{H_1\}$  then **S** asks the question ‘Is the secret  $H_1$ ?’ Suppose Nature responds ‘No’. This means that the only hypothesis that has been temporarily rejected is  $H_1$ . The information provided by this first round of the USI game can be expressed by the formula  $H_2 \vee H_3 \vee H_4$  and under the assumption that  $T$  holds, this is logically equivalent to  $\neg H_1$ . In general, if Nature’s answer is ‘Yes’ and  $I_D$  is the set of indexes of the hypotheses in  $D$ , then this information is expressed by the formula  $\bigvee_{i \in I_D} H_i$ , which, modulo  $T$ , is logically equivalent to  $\bigwedge_{i \in [n] \setminus I_D} \neg H_i$  ( $n = |\mathcal{H}|$ ). If Nature’s answer to the question ‘Does the secret belong to  $D \subseteq \mathcal{H}$ ?’ is ‘No’, then this information is expressed by the two logically equivalent formulas  $\bigvee_{i \in [n] \setminus I_D} H_i$  and  $\bigwedge_{i \in I_D} \neg H_i$ . Table 4 illustrates sample questions-and-answers relative to  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$ .

The *degree to which a hypothesis is temporarily rejected by a single question-and-answer*  $\delta$  in a USI game after  $k$  questions-and-answers, denoted  $r_{\delta}^u$ , is arrived at as follows. Before any question

TABLE 4. Formalization of data in a USI game

| Question                | Answer | Equivalent Formalizations (given $T$ ) |                                            |
|-------------------------|--------|----------------------------------------|--------------------------------------------|
| $D = \{H_1, H_2\}$      | Yes    | $H_1 \vee H_2$                         | $\neg H_3 \wedge \neg H_4$                 |
| $D = \{H_4\}$           | Yes    | $H_4$                                  | $\neg H_1 \wedge \neg H_2 \wedge \neg H_3$ |
| $D = \{H_1, H_3, H_4\}$ | No     | $H_2$                                  | $\neg H_1 \wedge \neg H_3 \wedge \neg H_4$ |
| $D = \{H_2\}$           | No     | $H_1 \vee H_3 \vee H_4$                | $\neg H_2$                                 |

is answered,  $r_\emptyset^u(H)$  equals 0 for each  $H \in \mathcal{H}$ . Then as data in the form of questions-and-answers arrive, let

$$r_\delta^u(H) = \begin{cases} \infty & \text{if } T \models \neg\delta, \\ 1 & \text{if } T, H \models \neg\delta, \\ 0 & \text{otherwise.} \end{cases} \quad (7)$$

As usual the aggregation of degrees of rejection yielded by multisets of data is additive, i.e. the sum of the degree of rejection of the formula occurring in the multiset.

There are however two exceptional cases where  $r_\Delta^u(H)$  is taken to be infinite, and different from the sum. The first arises when  $\sum_{\delta \in \Delta} r_\delta^u(H_{s^*}) > m$ . This accounts for the situation in which the secret has been rejected, i.e. temporarily rejected more times than the fixed number  $m$  of available lies. Since  $\mathbf{N}$  knows that the secret is true, this is a contradiction. (Recall that the blueprint consequence relation is explosive, see Lemma 2.4.) The second case in which  $r_\Delta^u(H_i)$  is infinite arises when hypotheses distinct from the secret are rejected, i.e.  $\sum_{\delta \in \Delta} r_\delta^u(H_i) > m$ .

Summing up, we let:

$$r_\Delta^u(H_i) = \infty \begin{cases} \text{if } H_i \in \mathcal{H} \text{ and } \sum_{\delta \in \Delta} r_\delta^u(H_{s^*}) > m \\ \text{or } H_i \neq H_{s^*} \text{ and } \sum_{\delta \in \Delta} r_\delta^u(H_i) > m \end{cases} \quad (8)$$

while, in the remaining cases:

$$r_\Delta^u(H_i) = \sum_{\delta \in \Delta} r_\delta^u(H_i). \quad (9)$$

Note that  $H_{s^*}$  might not even belong to  $\hat{\mathcal{H}}_\Delta$  for arbitrary  $\Delta$ . We are only guaranteed that  $r_\Delta^u(H_{s^*}) \leq m$ , i.e. it cannot be temporarily rejected more than  $m$  times.

#### LEMMA 4.2

$r^u$  is a degree of rejection.

PROOF. We need to check that the rejection degree defined in (8) and (9) satisfies Definition 2.1.

For condition (1), suppose  $\gamma \models \delta$ , for  $\gamma, \delta \in \text{Fm}$ . Since  $\gamma \models \delta$ , we have that  $\neg\delta \models \neg\gamma$ , hence if  $T \models \neg\delta$ , we have  $T \models \neg\gamma$  and, if  $T, H \models \neg\delta$  then  $T, H \models \neg\gamma$ . This, by equation (7) entails that  $r_\gamma^u(H) \geq r_\delta^u(H)$ .

The remaining conditions are immediate consequences of the definition, hence we omit checking them explicitly.  $\square$

#### EXAMPLE 4.3

Continuing with the setting of Example 4.1, let  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$ . Suppose now that the number of lies Nature can use is  $m = 2$ . Then, the following is a possible sequence of plays of the

USI game:

- (1) Question: 'Is the secret in  $D_1 = \{H_1, H_2\}$ ?'  
Answer: 'No'.
- (2) Question: 'Is the secret in  $D_2 = \{H_4\}$ ?'  
Answer: 'Yes'.
- (3) Question: 'Is the secret in  $D_3 = \{H_1, H_2, H_3\}$ ?'  
Answer: 'Yes'.
- (4) Question: 'Is the secret in  $D_4 = \{H_1, H_2, H_4\}$ ?'  
Answer: 'No'.
- (5) Question: 'Is the secret in  $D_5 = \{H_4\}$ ?'  
Answer: 'Yes'.
- (6) Question: 'Is the secret in  $D_6 = \{H_4\}$ ?'  
Answer: 'No'.

Denote with  $\delta_i$  the formalization of the  $i$ -th question-and-answer, i.e.

$$\begin{aligned}\delta_1 &= H_3 \vee H_4; \\ \delta_2 &= H_4; \\ \delta_3 &= H_1 \vee H_2 \vee H_3; \\ \delta_4 &= H_3; \\ \delta_5 &= H_4; \\ \delta_6 &= H_1 \vee H_2 \vee H_3.\end{aligned}$$

Thus, the six moves in the USI game are captured formally by the multiset

$$\Delta = \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3, H_3, H_4, H_1 \vee H_2 \vee H_3\}.$$

Let us denote with  $\Delta_i$  the multiset consisting of the formalizations of the first  $i$  questions-and-answers according to the sequence just introduced, i.e. in our specific example:

$$\begin{aligned}\Delta_1 &= \{H_3 \vee H_4\}; \\ \Delta_2 &= \{H_3 \vee H_4, H_4\}; \\ \Delta_3 &= \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3\}; \\ \Delta_4 &= \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3, H_3\}; \\ \Delta_5 &= \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3, H_3, H_4\}; \\ \Delta_6 &= \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3, H_3, H_4, H_1 \vee H_2 \vee H_3\}.\end{aligned}$$

Therefore, after each question-and-answer we can compute the degree of rejection for each hypotheses in  $\mathcal{H}$  according to (8) and (9). See Tables 5 and 6 for the computations of the degrees of rejection relative to this example.

Since  $r_{\Delta_6}^u(H_i) = \infty$  for  $i = 1, 2, 4$  and  $r_{\Delta_6}^u(H_3) = 2$ , i.e. all the hypotheses except for one have been rejected, it follows that  $H_3$  is the secret.

This example motivates our next Definition.

**DEFINITION 4.4** (uRJ-consequence).

Let  $\Delta$  be a multiset in  $\text{Fm}_{\mathcal{H}}$ ,  $\varphi \in \text{Fm}_{\mathcal{H}}$ , and  $r_{\Delta}^u(\varphi)$  be defined as in (8) and (9). Then define

$$\Delta \sim_{\text{uRJ}} \varphi \text{ if } T, H \models \varphi, \text{ for every } H \in \hat{\mathcal{H}}_{\Delta}.$$

Note that the fact we restrict data to answers in a USI game effectively means that uRJ-consequence relates hypotheses.

TABLE 5. Degrees of rejection computed on the single questions-and-answers  $\delta_i$ 

|                  | Hypotheses |       |       |       |
|------------------|------------|-------|-------|-------|
|                  | $H_1$      | $H_2$ | $H_3$ | $H_4$ |
| $r_{\delta_1}^u$ | 1          | 1     | 0     | 0     |
| $r_{\delta_2}^u$ | 1          | 1     | 1     | 0     |
| $r_{\delta_3}^u$ | 0          | 0     | 0     | 1     |
| $r_{\delta_4}^u$ | 1          | 1     | 0     | 1     |
| $r_{\delta_5}^u$ | 1          | 1     | 1     | 0     |
| $r_{\delta_6}^u$ | 0          | 0     | 0     | 1     |

TABLE 6. Degrees of rejection computed on the sequence of questions-and-answers  $\Delta_i$ 

|                  | Hypotheses |          |       |          |
|------------------|------------|----------|-------|----------|
|                  | $H_1$      | $H_2$    | $H_3$ | $H_4$    |
| $r_{\Delta_1}^u$ | 1          | 1        | 0     | 0        |
| $r_{\Delta_2}^u$ | 2          | 2        | 1     | 0        |
| $r_{\Delta_3}^u$ | 2          | 2        | 1     | 1        |
| $r_{\Delta_4}^u$ | $\infty$   | $\infty$ | 1     | 2        |
| $r_{\Delta_5}^u$ | $\infty$   | $\infty$ | 2     | 2        |
| $r_{\Delta_6}^u$ | $\infty$   | $\infty$ | 2     | $\infty$ |

A continuation of the Example 4.3 illustrates Definition 4.4.

#### EXAMPLE 4.5

Recall that  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$ , Nature can lie at most  $m = 2$  times, and the available data is formalized by the multiset  $\Delta = \{H_3 \vee H_4, H_4, H_1 \vee H_2 \vee H_3, H_3, H_4, H_1 \vee H_2 \vee H_3\}$ . For each  $\Delta_i$  each  $\hat{\mathcal{H}}_{\Delta_i}$  is computed as follows:

$$\hat{\mathcal{H}}_{\Delta_1} = \{H_3, H_4\};$$

$$\hat{\mathcal{H}}_{\Delta_2} = \{H_4\};$$

$$\hat{\mathcal{H}}_{\Delta_3} = \{H_3, H_4\};$$

$$\hat{\mathcal{H}}_{\Delta_4} = \{H_3\};$$

$$\hat{\mathcal{H}}_{\Delta_5} = \{H_3, H_4\};$$

$$\hat{\mathcal{H}}_{\Delta_6} = \{H_3\}.$$

Note that for every  $i, j$  s.t.  $i \neq j$  we have  $T, H_i \models \neg H_j$ . Therefore, the following consequence relations holds

$$\Delta_1 \vdash_{uRJ} \neg H_1 \wedge \neg H_2;$$

$$\Delta_2 \vdash_{uRJ} \neg H_3;$$

$$\Delta_6 \vdash_{uRJ} \neg H_1 \wedge \neg H_2 \wedge \neg H_4.$$

TABLE 7. Rules for mutually exclusive and exhaustive hypotheses

$$\begin{array}{c}
\frac{\Delta \vdash \varphi \quad T, \varphi \models \psi}{\Delta \vdash \psi} \text{ (RWE)} \qquad \frac{\Delta, \gamma \vdash \varphi \quad T \models \gamma \leftrightarrow \delta}{\Delta, \delta \vdash \varphi} \text{ (LLE)} \\
\\
\frac{\Delta, \gamma, \delta \vdash \varphi \quad T, \neg \gamma, \neg \delta \models \perp}{\Delta, \gamma \wedge \delta \vdash \varphi} \text{ (AND)}_I \quad \frac{\Delta, \gamma \wedge \delta \vdash \varphi \quad T, \neg \gamma, \neg \delta \models \perp}{\Delta, \gamma, \delta \vdash \varphi} \text{ (AND)}_I^{con}
\end{array}$$

On the other hand, we have

$$\begin{aligned}
\Delta_1 &\not\sim_{uRJ} \neg H_3; \\
\Delta_1 &\not\sim_{uRJ} \neg H_4; \\
\Delta_2 &\not\sim_{uRJ} \neg H_4; \\
\Delta_6 &\not\sim_{uRJ} \neg H_3.
\end{aligned}$$

Our next lemma collects immediate consequences of Definition 4.4, which will be useful later on.

#### LEMMA 4.6

Let  $\mathcal{H} = \{H_1, \dots, H_n\}$  and let  $m \geq 0$  be the number of lies allowed to Nature. The following hold:

- (1) Suppose that  $\hat{\mathcal{H}}_\Delta = \hat{\mathcal{H}}_{\Delta'}$ , then  $\Delta \sim_{uRJ} \varphi$  if and only if  $\Delta' \sim_{uRJ} \varphi$ .
- (2) If  $\Delta$  rejects  $H_i$ , then  $\Delta \vdash_{uRJ} \neg H_i$ .
- (3) If  $\Delta$  rejects  $H_i$ , then  $\Delta' \vdash_{uRJ} \neg H_i$  for every  $\Delta' \supseteq \Delta$ .
- (4)  $H_{s^*}$  is the secret if and only if for some  $\Delta \in \text{Fm}$  consistent with  $T$ , we have  $\hat{\mathcal{H}}_\Delta = \{H_{s^*}\}$  and  $r_\Delta^u(H_i) = \infty$  for all  $H_i \neq H_{s^*}$ .

To show that indeed  $\sim_{uRJ}$  arises from the blueprint  $\sim_{RJ}$  consequence relation we need to formulate a slightly stronger version of the rules (RWE) and (LLE) where the theory  $T$  is mentioned explicitly. (Recall  $T$  expresses that the hypotheses form a partition containing the unique *secret*). The resulting amendments are presented in Table 7, which also contains the strengthening of  $(\text{AND})_I$  and  $(\text{AND})_I^{con}$ . Note that, with a slight abuse of notation, we do not change the name of the rules compared to Table 2.

As an immediate consequence of Proposition 2.6 above, both (AMON) and (MMON) fail for  $\sim_{uRJ}$ , as expected. Moreover a straightforward adaptation of the arguments provided by the proof of Lemma 2.5 and Lemma 3.3 will yield the following Lemma.

#### LEMMA 4.7

$\sim_{uRJ}$  satisfies the rules in Table 7.

In light of this we can observe that  $\sim_{uRJ}$  arises from  $\sim_{RJ}$  by taking  $r = r^u$ . To simplify notation in the remainder of this Section we will just write  $r$  and speak of the *degree of rejection*.

### 4.2 Valid rules of inference for $\sim_{uRJ}$

Since its language allows for a formula to be both on the left and on the right of the consequence relation symbol, we can ask whether  $\sim_{uRJ}$  satisfies (multiset versions of) (CUT), (CMO) and (RMO) described in Table 8. Quite interestingly,  $\sim_{uRJ}$  also satisfies a form of constrained monotonicity,

TABLE 8. Rules satisfied by the System  $\vdash_{uRJ}$ 

$$\begin{array}{c}
\frac{}{\varphi \vdash_{RJ} \varphi} \text{ (REF)} \\
\\
\frac{\Delta \vdash_{uRJ} \psi \quad \Delta \vdash_{uRJ} \varphi}{\Delta, \varphi \vdash_{uRJ} \psi} \text{ (CMO)} \quad \frac{\Delta, \varphi \vdash_{RJ} \psi \quad \Delta \vdash_{RJ} \varphi}{\Delta \vdash_{RJ} \psi} \text{ (CUT)} \\
\\
\frac{\Delta \vdash \psi \quad \Delta \not\vdash \neg \varphi}{\Delta, \varphi \vdash \psi} \text{ (RMO)} \quad \frac{\Delta \vdash \psi \quad \{\Delta \setminus \{\delta\} \vdash \psi\}_{\delta \in \Delta}}{\Delta, \gamma \vdash \psi} \text{ (UMO)}
\end{array}$$

which to the best of our knowledge has not been investigated before, and which we term (UMO), for USI games:

$$\frac{\Delta \vdash \psi \quad \{\Delta \setminus \{\delta\} \vdash \psi\}_{\delta \in \Delta}}{\Delta, \varphi \vdash \psi} \text{ (UMO)}$$

Before proving its validity, we illustrate the rationale for (UMO). The non-monotonic behaviour of  $\vdash_{uRJ}$  is ultimately due to the fact that new questions-and-answers can change the set of least rejected hypotheses. Note however that a *single* question-and-answer may alter the degree of rejection of each hypothesis by at most by one. Hence, when the degree of rejection of any least rejected hypothesis, say  $H$ , is at least two units smaller than the others in  $\mathcal{H}$ , we can be sure that adding just one question-and-answer, say  $\varphi$ , will not change the status of  $H$  as a least rejected hypothesis. This condition on the ‘distance’ between degrees of rejection of least rejected hypotheses and the others is captured by the second premise of (UMO). If the degree of rejection of the least rejected hypotheses are smaller by at the least two units than the others in  $\mathcal{H}$ , then removing any occurrence of a formula from the multiset  $\Delta$  will not change  $\hat{\mathcal{H}}_\Delta$ , and hence the formulas that the multiset entails.

#### EXAMPLE 4.8

Suppose

$$H_1, H_1 \vdash_{uRJ} H_1.$$

Here  $\hat{\mathcal{H}}_{\{H_1, H_1\}} = \{H_1\}$ ,  $r_{\{H_1, H_1\}}(H_1) = 0$  and  $r_{\{H_1, H_1\}}(H_i) = 2$  for any  $H_i \in \mathcal{H}$ , with  $i \neq 1$ . This ensures that

$$H_1, H_1, \varphi \vdash_{uRJ} H_1$$

for any formula  $\varphi$ , since we may have, at worst,  $r_{\{H_1, H_1, \varphi\}}(H_1) = 1$ , while  $r_{\{H_1, H_1, \varphi\}}(H_i) \geq 2$  for any  $i \neq 1$ .  $H_1$  will therefore continue to be the only formula in  $\hat{\mathcal{H}}_{\{H_1, H_1, \varphi\}}$ . The above reasoning is thus captured by the following application of (UMO):

$$\frac{H_1, H_1 \vdash_{uRJ} H_1 \quad H_1 \vdash_{uRJ} H_1}{H_1, H_1, H_2 \vdash_{uRJ} H_1}$$

For a non-example, note that

$$\frac{H_1 \vdash_{uRJ} H_1}{H_1, H_2 \vdash_{uRJ} H_1}$$

is not valid. To have a correct instance of (UMO), we should also have  $\vdash_{uRJ} H_1$ , which clearly does not hold. In the above consequence we have  $H_1 \in \hat{\mathcal{H}}_{H_1}$ , but the difference between  $r_{H_1}(H_2) = 1$  and  $r_{H_1}(H_1) = 0$  is not strictly greater than 1. In the consequence we now have  $r_{\{H_1, H_2\}}(H_1) = r_{\{H_1, H_2\}}(H_2) = 1$ , and both  $H_1$  and  $H_2$  belong to  $\hat{\mathcal{H}}_{\{H_1, H_2\}}$ , hence  $H_1, H_2 \not\vdash_{uRJ} H_1$ .

PROPOSITION 4.9

$\vdash_{uRJ}$  satisfies the rules in Table 8.

Thus  $\vdash_{uRJ}$  is ‘more monotonic’ than rational consequence relations. So it captures a rather minimal departure from classical deduction in which nonmonotonic behaviour arises as a consequence of lies in the USI game, which corresponds to a variety of errors that scientists can make in experimental research. This vicinity with classical deduction can be interpreted as data-driven reasoning being a good approximation of the gold standard of mathematical proof.

#### 4.3 Disjunction in the premisses and Modus tollens hold for $\vdash_{uRJ}$

In addition to satisfying (UMO) and (RMO),  $\vdash_{uRJ}$  goes beyond  $\vdash_{IRJ}$  by satisfying also the (OR) rule. In light of our comments on Lemma 3.4 this welcome feature of  $\vdash_{uRJ}$  should not come as a surprise upon careful comparison of (5) and (7). While the semantics governing  $\vdash_{IRJ}$  is quantitative (i.e. the full range of probability functions plays a role in (5)),  $\vdash_{uRJ}$  is qualitative (i.e. only the extreme points of  $[0, 1]$  play any role in (7)).

LEMMA 4.10

$\vdash_{uRJ}$  satisfies

$$\frac{\Delta, \gamma \vdash \psi \quad \Delta, \delta \vdash \psi}{\Delta, \gamma \vee \delta \vdash \psi.} \text{ (OR)}$$

See the [Appendix](#) for the proof.

The qualitative nature of the degree of rejection underpinning  $\vdash_{uRJ}$  gives us also a version of *modus tollens*. This conspicuous rule, which to some is the key pattern in scientific reasoning, typically fails for quantitative inference. A fact that has not gone unnoticed to some critics of NHST, as we recalled in Section 1.1 above. Hence the following Lemma, proven in the [Appendix](#), is of conceptual interest.

LEMMA 4.11

$\vdash_{uRJ}$  satisfies

$$\frac{\Delta, \varphi \vdash_{uRJ} \psi \quad \Delta \vdash_{uRJ} \neg \psi}{\Delta \vdash_{uRJ} \neg \varphi.} \text{ (MT)}$$

We now turn to the main result of this paper.

#### 4.4 A completeness result for $\vdash_{uRJ}$

We make our way now to establishing a completeness result for  $\vdash_{uRJ}$ -consequence relations. In preparation for that, the next Lemma pins down a useful normal form for arbitrary consequences under  $\vdash_{uRJ}$ . We will show that whenever  $\vdash_{uRJ}$  relates a multiset of premisses and a conclusion, we

## 22 A logical framework for data-driven reasoning

can find an equivalent set of consequence relations  $\sim_{uRJ}$ , with ‘simpler’ premises and conclusion. This will be useful in establishing the argument by cases, which intervenes in the proof of Theorem 4.14.

LEMMA 4.12

Let  $\Delta$  be a multiset of formulas in  $\text{Fm}$  and  $\psi \in \text{Fm}$ . The following are equivalent:

- (1)  $\Delta \sim_{uRJ} \psi$ ;
- (2) A set of consequence relations of the form  $\neg H_1^{l_1}, \dots, \neg H_n^{l_n} \sim_{uRJ} \neg H_j$  holds, with multiplicity indices  $l_1, \dots, l_n$  greater or equal than 0 and  $j \in \{1, \dots, n\}$ .

PROOF. (1)  $\Rightarrow$  (2). Assume  $\Delta \sim_{uRJ} \psi$ . We proceed as follows. First, note that any formula  $\varphi$  in  $\text{Fm}_{\mathcal{H}}$  is logically equivalent, modulo  $T$ , to the disjunction of the hypotheses entailing it, i.e. to the formula

$$\varphi_{\vee} := \bigvee_{H \models \varphi} H.$$

In turn, since under our assumptions  $H \not\models \varphi$  entails  $H \models \neg\varphi$ , the formula  $\varphi_{\vee}$  is equivalent to:

$$\varphi_{\wedge} := \bigwedge_{H \models \neg\varphi} \neg H.$$

In the limiting case where  $\{H \in \mathcal{H} \mid H \models \varphi\} = \mathcal{H}$ , and consequently  $\{H \in \mathcal{H} \mid H \models \neg\varphi\} = \emptyset$ , we let  $\varphi_{\vee} \equiv \varphi_{\wedge} \equiv T$ . We then let  $\Delta_{\wedge}$  be the multiset obtained by replacing each  $\delta \in \Delta$  by  $\delta_{\wedge}$ . Since all the formulas in  $\Delta_{\wedge}$  and  $\psi$  are logically equivalent (modulo  $T$ ) to  $\Delta$  and  $\psi$ , we will then have that  $\Delta \sim_{uRJ} \psi$  if and only if

$$\Delta_{\wedge} \sim_{uRJ} \psi_{\wedge}.$$

From the latter, we can derive the original  $\Delta \sim_{uRJ} \psi$  by repeated backwards applications of (LLE) to recover  $\Delta$ , and (RWE) to recover  $\psi$

Now, if one of the formulas in  $\Delta$ , say  $\delta_{\wedge}$ , is equivalent to  $T$ , it does not temporarily reject any noncontradictory formula, hence

$$\Delta_{\wedge} \sim_{uRJ} \psi_{\wedge} \text{ if and only if } \Delta_{\wedge} \setminus \{T\} \sim_{uRJ} \psi_{\wedge}.$$

From it we can derive  $\Delta_{\wedge} \sim_{uRJ} \psi_{\wedge}$  by just applying (RMO), since  $\Delta_{\wedge} \setminus \{T\} \not\sim_{uRJ} \neg T$ , i.e.

$$\frac{\Delta_{\wedge} \setminus \{T\} \sim_{uRJ} \psi_{\wedge} \quad \Delta_{\wedge} \setminus \{T\} \not\sim_{uRJ} \neg T}{\Delta_{\wedge} \sim_{uRJ} \psi_{\wedge}} \text{ (RMO)}$$

We may thus remove from  $\Delta$  all the formulas that are logically equivalent to  $T$ .

Our next step is applying backwards the rules (AND)<sub>I</sub> and (AND). Note that we can equivalently take the converse (AND)<sub>I</sub><sup>con</sup>, since the conjuncts are of the form  $\neg H_1 \wedge \neg H_2 \wedge \dots$ , and their negations  $H_1, H_2, \dots$  are all mutually exclusive modulo  $T$ , hence they satisfy the condition for applying (AND)<sub>I</sub><sup>con</sup>. We thus reduce  $\Delta_{\wedge} \sim_{uRJ} \psi_{\wedge}$  to the set of consequences of the form:

$$\neg H_1^{l_1}, \dots, \neg H_n^{l_n} \sim_{uRJ} \neg H_j$$

as required.

(2)  $\Rightarrow$  (1). The result follows by applying (forwards, this time) all the rules used in the previous step.  $\square$

The normal form depends on the set of hypotheses  $\mathcal{H}$ , as the next example illustrates.

## EXAMPLE 4.13

Consider first the normal form of  $H_1 \vee H_2, H_1 \vee H_3 \vdash_{uRJ} H_1$  assuming  $\mathcal{H} = \{H_1, H_2, H_3\}$ :

$$\delta_1 := H_1 \vee H_2$$

$$\delta_2 := H_1 \vee H_3$$

$$\psi := H_1$$

$$\delta_{1\wedge} := \neg H_3$$

$$\delta_{2\wedge} := \neg H_2$$

$$\psi_{\wedge} := \neg H_2 \wedge \neg H_3$$

Thus,  $\Delta_{\wedge} = \{\neg H_3, \neg H_2\}$  and  $\hat{\mathcal{H}}_{\Delta_{\wedge}} = \{H_1\}$ .

Note that if we let  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$  we obtain the following:

$$\delta'_{1\wedge} := \neg H_3 \wedge \neg H_4$$

$$\delta'_{2\wedge} := \neg H_2 \wedge \neg H_4$$

$$\psi'_{\wedge} := \neg H_2 \wedge \neg H_3 \wedge \neg H_4$$

Hence  $\Delta'_{\wedge} = \{\neg H_3, \neg H_2, \neg H_4\}$  and  $\hat{\mathcal{H}}_{\Delta'_{\wedge}} = \{H_1\}$ .

It is easy to see that, in case  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$  we have  $\Delta'_{\wedge} \not\vdash_{uRJ} H_1$ .

Now to our main result.

## THEOREM 4.14

Suppose  $\Delta \vdash_{uRJ} \psi$  and  $r_{\Delta}(H)$  is finite for each  $H \in \mathcal{H}$ . Then there is a derivation of it using the rules in Table 7 and Table 8.

The proof is in the [Appendix](#). Here we sketch the idea. Lemma 4.12 guarantees the existence of a set of

$$\neg H_1^{l_1}, \dots, \neg H_n^{l_n} \vdash_{uRJ} \neg H_j,$$

where all the multiplicity indices  $l_1, \dots, l_n$  are greater or equal than 0 and  $j \in \{1, \dots, n\}$ . Pick any such consequences and denote its left-hand side by  $\Delta'$ . The proof then proceeds reasoning by induction on the number of formulas occurring in  $\Delta'$ . We need to show how to derive  $\Delta' \vdash_{uRJ} \neg H_j$  by a rule in either Table 7 or Table 8, using valid premises with a number of formulas smaller than  $\Delta'$ . Example 4.15 below illustrates some of the relevant subcases, involving (UMO) and (RMO).

## EXAMPLE 4.15

We illustrate by way of examples how to apply the rules (UMO) and (RMO) in the corresponding sub-case of the proof of Theorem 4.14. In each case below, we set  $\mathcal{H} = \{H_1, H_2, H_3, H_4\}$ .

- (1) Suppose that  $\Delta' = \{\neg H_4^3\}$  where 3 is the multiplicity index relative to  $\neg H_4$ . Thus,  $\hat{\mathcal{H}}_{\Delta'} = \mathcal{H} \setminus \{H_4\} = \{H_1, H_2, H_3\}$ , and  $\neg H_4^3 \vdash_{uRJ} \neg H_4$  holds, while  $H_i \models \neg H_4$ , for every  $i = 1, 2, 3$ . If we remove  $\neg H_4$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_4\} = \{\neg H_4^2\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_4\}} = \{H_1, H_2, H_3\}$ . Therefore,  $\Delta' \setminus \{\neg H_4\} \vdash_{uRJ} \neg H_4$ . In addition, for every  $i = 1, 2, 3$  we have  $H_i \not\models H_4$  and  $\Delta' \setminus \{\neg H_4\} \not\vdash_{uRJ} H_4$ .

In conclusion  $\neg H_4^3 \vdash_{uRJ} \neg H_4$  can be obtained by applying (RMO) to  $\neg H_4^2 \vdash_{uRJ} \neg H_4$  as first premise and  $\neg H_4^2 \not\vdash_{uRJ} H_4$  as second premise, i.e.

$$\frac{\neg H_4^2 \vdash_{uRJ} \neg H_4 \quad \neg H_4^2 \not\vdash_{uRJ} H_4}{\neg H_4^3 \vdash_{uRJ} \neg H_4} \text{ (RMO)}.$$

Suppose that  $\Delta' = \{\neg H_1^2, \neg H_4^3\}$ . Thus,  $\hat{\mathcal{H}}_{\Delta'} = \{H_2, H_3\}$ . Since  $H_2 \models \neg H_1$  and  $H_3 \models \neg H_1$ , then  $\neg H_1^2, \neg H_4^3 \vdash_{uRJ} \neg H_1$  holds. If we remove  $\neg H_4$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_4\} = \{\neg H_1^2, \neg H_4^2\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_4\}} = \{H_2, H_3\}$ . Therefore,  $\Delta' \setminus \{\neg H_4\} \vdash_{uRJ} \neg H_4$ . In addition, for every  $i = 2, 3$   $H_i \not\models H_1$  and  $\Delta' \setminus \{\neg H_4\} \not\vdash_{uRJ} H_1$ .

In conclusion  $\neg H_1^2, \neg H_4^3 \vdash_{uRJ} \neg H_1$  can be obtained by applying (RMO) to  $\neg H_1^2, \neg H_4^2 \vdash_{uRJ} \neg H_1$  as first premise and  $\neg H_1^2, \neg H_4^2 \not\vdash_{uRJ} H_4$  as second premise, i.e.

$$\frac{\neg H_1^2, \neg H_4^2 \vdash_{uRJ} \neg H_1 \quad \neg H_1^2, \neg H_4^2 \not\vdash_{uRJ} H_4}{\neg H_1^2, \neg H_4^3 \vdash_{uRJ} \neg H_1} \text{ (RMO)}.$$

- (2a) Suppose that  $\Delta' = \{\neg H_1^2, \neg H_2^1, \neg H_3^3, \neg H_4^3\}$ . Thus,  $\hat{\mathcal{H}}_{\Delta'} = \{H_2\}$ . Since  $H_2 \models \neg H_4$ , then  $\neg H_1^2, \neg H_2^1, \neg H_3^3, \neg H_4^3 \vdash_{uRJ} \neg H_4$  holds. If we remove  $\neg H_3$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_3\} = \{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_3\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_3\} \vdash_{uRJ} \neg H_4$ . In addition,  $H_2 \not\models H_3$  and  $\Delta' \setminus \{\neg H_3\} \not\vdash_{uRJ} H_3$ .

In conclusion  $\neg H_1^2, \neg H_2^1, \neg H_3^3, \neg H_4^3 \vdash_{uRJ} \neg H_4$  can be obtained by applying (RMO) to  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4$  as first premise and  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \not\vdash_{uRJ} H_3$  as second premise, i.e.

$$\frac{\neg H_1^2, \neg H_2^1, \neg H_3^3, \neg H_4^3 \vdash_{uRJ} \neg H_4 \quad \neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \not\vdash_{uRJ} H_3}{\neg H_1^2, \neg H_2^1, \neg H_3^3, \neg H_4^3 \vdash_{uRJ} \neg H_4} \text{ (RMO)}.$$

- (2a') Suppose that  $\Delta' = \{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3\}$ . Thus,  $\hat{\mathcal{H}}_{\Delta'} = \{H_2\}$ . Since  $H_2 \models \neg H_4$ , then  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4$  holds.

If we remove  $\neg H_2$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_2\} = \{\neg H_1^2, \neg H_3^2, \neg H_4^3\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_2\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_2\} \vdash_{uRJ} \neg H_4$ .

If we remove  $\neg H_2$  and  $\neg H_1$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_2, \neg H_1\} = \{\neg H_3^2, \neg H_4^3\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_2, \neg H_1\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_2, \neg H_1\} \vdash_{uRJ} \neg H_4$ .

If we remove  $\neg H_2^2$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_2^2\} = \{\neg H_1^2, \neg H_3^2, \neg H_4^3\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_2^2\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_2^2\} \vdash_{uRJ} \neg H_4$ .

If we remove  $\neg H_2$  and  $\neg H_3$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_2, \neg H_3\} = \{\neg H_1^2, \neg H_3, \neg H_4^3\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_2, \neg H_3\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_2, \neg H_3\} \vdash_{uRJ} \neg H_4$ .

If we remove  $\neg H_2$  and  $\neg H_4$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_2, \neg H_4\} = \{\neg H_1^2, \neg H_3^2, \neg H_4^2\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_2, \neg H_4\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_2, \neg H_4\} \vdash_{uRJ} \neg H_4$ .

In conclusion  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4$  can be obtained by applying (UMO) to  $\neg H_1^2, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4$  as first premise and  $\{\{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3\} \setminus \{\neg H_2, \neg H_4\} \vdash_{uRJ} \neg H_4\} \neg H \in \Delta'$ , i.e.

$$\frac{\neg H_1^2, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4 \quad \{\{\neg H_1^2, \neg H_3^2, \neg H_4^3\} \setminus \{\neg H_2\} \vdash_{uRJ} \neg H_4\} \neg H \in \Delta'}{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_4} \text{ (UMO)}.$$

- (2b) Suppose that  $\Delta' = \{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3\}$ . Thus,  $\hat{\mathcal{H}}_{\Delta'} = \{H_2\}$ . Since  $H_2 \models \neg H_1$ , then  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_1$  holds. If we remove  $\neg H_4$  from  $\Delta'$ , then  $\Delta' \setminus \{\neg H_4\} = \{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^2\}$  and  $\hat{\mathcal{H}}_{\Delta'} = \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_4\}} = \{H_2\}$ . Therefore,  $\Delta' \setminus \{\neg H_4\} \vdash_{uRJ} \neg H_4$ . In addition,  $H_2 \not\models H_1$  and  $\Delta' \setminus \{\neg H_4\} \not\vdash_{uRJ} H_1$ .

TABLE 9. A comparison of relevant properties satisfied by rejection-based consequence relations

|                      | AND <sub>I</sub> | AND <sub>I</sub> <sup>con</sup> | OR | UMO |
|----------------------|------------------|---------------------------------|----|-----|
| $\vdash_{IRJ}$       | ✓                | ✓                               |    |     |
| $\vdash_{\alpha RJ}$ |                  |                                 |    |     |
| $\vdash_{uRJ}$       | ✓                | ✓                               | ✓  | ✓   |

In conclusion  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_1$  can be obtained by applying (RMO) to  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_1$  as first premise and  $\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \not\vdash_{uRJ} H_4$  as second premise, i.e.

$$\frac{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_1 \quad \neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \not\vdash_{uRJ} H_4}{\neg H_1^2, \neg H_2^1, \neg H_3^2, \neg H_4^3 \vdash_{uRJ} \neg H_1} \text{ (RMO)}.$$

#### EXAMPLE 4.16

We further illustrate Theorem 4.14 by showing how to find a derivation for the normal forms identified in Example 4.13. To this end, let  $\mathcal{H} = \{H_1, H_2, H_3\}$  and  $\neg H_2, \neg H_3 \vdash_{uRJ} \neg H_2 \wedge \neg H_3$ . Thus,

$$\frac{\frac{\neg H_2 \vdash_{uRJ} \neg H_2 \quad \neg H_2 \not\vdash_{uRJ} H_3}{\neg H_2, \neg H_3 \vdash_{uRJ} \neg H_2} \text{ (RMO)} \quad \frac{\neg H_3 \vdash_{uRJ} \neg H_3 \quad \neg H_3 \not\vdash_{uRJ} H_2}{\neg H_2, \neg H_3 \vdash_{uRJ} \neg H_3} \text{ (RMO)}}{\neg H_2, \neg H_3 \vdash_{uRJ} \neg H_2 \wedge \neg H_3} \text{ (AND)},$$

is the required derivation.

## 5 Conclusion and further work

We have introduced a family of consequence relations with the following intended semantics:  $\varphi$  follows from  $\Delta$  if  $\varphi$  is a classical consequence of every hypothesis that is least rejected in  $\Delta$ . The formalization of this leads to our blueprint consequence relation of Definition 2.3. This in turn produces three distinct, but related consequence relations:  $\vdash_{IRJ}$ ,  $\vdash_{\alpha RJ}$  and  $\vdash_{uRJ}$ . The first,  $\vdash_{IRJ}$ , is based on a rejection function that pins down maximal likelihood hypotheses. The rejection function underpinning  $\vdash_{\alpha RJ}$  is based on a logical rendering of the p-value. Finally, the rejection function defining  $\vdash_{uRJ}$  arises by counting of how many times a hypothesis is temporarily rejected in a USI game.

We showed that  $\vdash_{IRJ}$ ,  $\vdash_{\alpha RJ}$  and  $\vdash_{uRJ}$  are variations on the theme of rational consequence relations, as is their precursor: the blueprint consequence relation  $\vdash_{RJ}$ . Table 9 provides a comparison (we omit from the table those rules that are satisfied by all of them, i.e. REF, LLE, RWE, AND, CUT, RMO, CMO).

The main result of the paper, Theorem 4.14, identifies the rules of inference with respect to which  $\vdash_{uRJ}$  is complete.

While laying down the intended semantics of our consequence relations, we detailed why we considered the rejection degrees desirable for maximum likelihood, null hypothesis significance tests and strong inference, respectively. To us this lends support to the blueprint consequence relation

being a promising starting point for a logical investigation into the general properties of data-driven inference, independently of one's own preferred view on the nature of probability.

Given the enormous variety and complexity of scientific inference, what we have put forward here is just a preliminary set of results. We hope that this is enough to encourage more logicians to take up the challenge. Not only will this help revamping the decidedly outdated picture scientists have of logic, but it may also contribute to bringing some of the most heated debates on statistical significance and, more generally, on the methodology of data-driven inference on more logical and less ideological grounds.

### 5.1 *Related and future work*

In a series of works culminating in [33], Henry Kyburg and Choh Man Teng investigate the problem of putting statistical inference on a non-monotonic logical footing. They do so by taking as a starting point the defeasible nature of rejecting  $H_0$  in NHST. To the best of our knowledge they are among the very first ones to do this from a formal-logical point of view. It is therefore appropriate to describe briefly how the present work departs from the pioneering contribution of Kyburg and Teng's.

Conceptually, [33] focusses on the fact that in NHST and, more generally in statistical inference, one usually justifies conclusions (about the population) based on the representativeness of the sample. So its aim, as far as the foundations of statistics is concerned, is to contribute to the long-standing *reference class problem*. Our motivating question is broader, as we pointed out in the introductory section. Specifically, while we acknowledge the paramount importance of statistical data in scientific inference, we do not assume that it is the only kind of data relevant to pin down the logical properties of scientific inference. This is why we encompass the more general problem of *strong inference*, which we capture through USI games. Our second, and more formal, point of departure from Kyburg and Teng is the specific non-monotonic framework. Given their research question, Default Logic [49] is a most natural setup, which they then expand to investigate the role of probabilistic evidence in providing a semantics for the justification of non-normal defaults. More precisely, the authors insist that the justification for the application of a default rule encodes the lack of evidence to the effect that the statistical sample used in the inference is atypical or biased. In line with the broader question asked, our results are formulated in terms of non-monotonic consequence relations of the preferential kind. As is well known [36] the two frameworks are related, but distinct. Similarly, and more importantly, we do not commit to any specific view on the meaning of probability, which is interpreted evidentially by Kyburg and Teng. Since their uncertainty representation belongs to the wider class of imprecise probability, we are at present unable to carry out a direct comparison with their results. This, however, we set out to do in future work.

One problem that arises in the context of rejection-based scientific inference is handling conflicting evidence. Kyburg and Teng tackle it at the level of the measure of uncertainty, which for them is evidential probability as just recalled. However, it is of technical and conceptual interest to focus on the fact that rejection-based inference can be seen as a form of paraconsistent inference. Recall that according to Fisher [16]:

[t]he interest of statistical tests for scientific workers depends entirely from their use in rejecting hypotheses which are thereby judged to be incompatible with the observations.

As captured by USI games, if the observations are classically inconsistent with a set of hypotheses, at least one of them must be rejected. But on more fine-grained analysis, this kind of 'incompatibility' will come in degrees. While the consequence relation  $\vdash_{\alpha RJ}$  tackles this by mimicking the calculated p-value, it is interesting to address the 'degree of incompatibility'

interpretation of p-values through a paraconsistent consequence relation, for which a vast landscape of logics are available [7, 9]. As a first line of work in this direction, it appears promising to distinguish *evidence in favour* from *evidence against* a given hypothesis. In this way a formal link with First-Degree Entailment [43] becomes available. This would then lead to revising our blueprint consequence relation to encompass logics of formal inconsistency [8, 10]. In particular, the probability functions defined over such systems [11, 29] are expected to give rise to rejection degrees capable of representing useful inference in light of conflicting evidence, which is very much the norm in scientific practice. The statistical literature on this is already quite promising [5, 14, 27, 53], see also [45] for an overview.

This brings us, in conclusion, to a logical take on statistical inference that has been put forward in the framework of dynamic doxastic logics by Baltag, Rafiee Rad and Smets (BRS) in [3]. In it, the authors model an agent performing statistical inference, as combining ‘strong’, epistemic information about the probability of an event, with ‘soft’, doxastic, information. The latter is construed as a plausibility function, which determines an ordering over the epistemically permissible probability distributions, reflecting the subjective inclination of the agent and the non-definitive evidence she is confronted with. The agent is then taken to believe only in the *maximal* distributions, according to her plausibilistic order. While the research question underlying the BRS framework is quite distinct from ours, a foundational commonality with the present work stands out: inference arises against some underlying epistemic ordering and background information. This is not surprising, since some epistemic ordering is always at the root of a significance test.

In conclusion, let us restate the key message of our results: consequence relations based on degrees of rejection are promising in the investigation of the validity of data-driven inference. A clear difficulty in making the next step, and ascertain whether they will also deliver in real-world cases of scientific inference, is to do with the fact that these are only partially formal. By this we mean that unlike mathematical proofs, arguments in data-driven science depend, to some hard-to-pin-down extent to the actual content of the specific reasoning at hand. This is why we are working on a bottom-up approach where specific instances of data-driven inference that a relevant community takes to be valid is represented within our system of consequence relations and hopefully abstracted to show its general, yet context-dependent, validity. Key tools in doing this are the further specialization of the rejection functions, and in particular alternatives to addition in the aggregation of the rejection offered by single pieces of data, and a more fine grained interpretation of the meaning of *lies* in the USI game. We hope to report encouraging results in this direction in future work.

## Acknowledgements

We are very grateful to the organizers and audiences of the LIRa seminar in Amsterdam, the IIIA-CSIC seminar in Barcelona, the triennial international conference of the Italian Society for Logic and the Philosophy of Science held at the University of Urbino, the MOSAIC Workshop held in Vienna.

Paolo Baldi acknowledges funding from the research project ‘Green-Aware MEchanismsS’ (GAMES), part of the PNRR MIUR project FAIR—Future AI Research (PE00000013), Spoke 9—Green-Aware AI.

Hykel Hosni acknowledges funding from the Italian Ministry for University and Research under the scheme FIS1—Reasoning with Data (ReDa) G53C23000510001.

Corsi and Hosni acknowledge funding by the Department of Philosophy ‘Piero Martinetti’ of the University of Milan under the Project ‘Departments of Excellence 2023–2027’ awarded by the

Italian Ministry of Education, University and Research (MIUR), and by the MOSAIC project (EU H2020-MSCA-RISE-2020 Project 101007627).

## Appendix

### PROPOSITION A.1

RJ-consequence satisfies the rules in Table 2.

PROOF. The proof proceeds by considering each rule in turn.

For (RWE), assume that  $\Delta \vdash_{RJ} \delta$  and  $\delta \models \psi$ . Since  $\Delta \vdash_{RJ} \delta$  for any  $H \in \hat{\mathcal{H}}_\Delta$  we have that  $T, H \models \delta$ . From  $\delta \models \psi$ , we then have, for any  $H \in \hat{\mathcal{H}}_\Delta$ , that  $T, H \models \psi$ .

(LLE) holds, since if we take two formulas  $\gamma$  and  $\delta$  that are logically equivalent, by (2) in Definition 2.1 we have  $r_\gamma(H) = r_\delta(H)$  for any  $H \in \mathcal{H}$ . Hence  $\hat{\mathcal{H}}_{\Delta, \gamma} = \hat{\mathcal{H}}_{\Delta, \delta}$  and  $\Delta, \gamma \vdash_{RJ} \varphi$  entails  $\Delta, \delta \vdash_{RJ} \varphi$ .

We now show that (AND) holds. Indeed, assume that  $\Delta \vdash_{RJ} \varphi$  and  $\Delta \vdash_{RJ} \psi$ . For any  $H \in \hat{\mathcal{H}}_\Delta$ , we then have that  $T, H \models \varphi$  and  $T, H \models \psi$ . Hence  $T, H \models \varphi \wedge \psi$  for any  $H \in \hat{\mathcal{H}}_\Delta$ .  $\square$

### PROPOSITION A.2

Neither (AMON) nor (MMON) hold for  $\vdash_{RJ}$ .

PROOF. For (AMON), let  $\mathcal{H} = \{H_1, H_2\}$  and take  $\gamma \vdash_{RJ} H_1$  to be the premise of (AMON), assuming that  $r_\gamma(H_1) \leq r_\gamma(H_2)$ . Now, we will have  $r_{\gamma \wedge \delta}(H_1) \geq r_\gamma(H_1)$  and  $r_{\gamma \wedge \delta}(H_2) \geq r_\gamma(H_2)$  by the properties of rejection functions. However, we may still have  $r_{\gamma \wedge \delta}(H_1) > r_{\gamma \wedge \delta}(H_2)$ , i.e.  $\hat{\mathcal{H}}_{\gamma \wedge \delta} = \{H_2\}$ , hence  $\gamma \wedge \delta \not\vdash_{RJ} H_1$ .

As for (MMON), assume again  $\mathcal{H} = \{H_1, H_2\}$  and  $\gamma \vdash_{RJ} H_1$  with  $r_\gamma(H_1) = 0$  and  $r_\gamma(H_2) = 1$ . Let  $\delta$  be such that  $r_\delta(H_1) = 1$  and  $r_\delta(H_2) = 0$ . We will have  $r_{\gamma, \delta}(H_1) = r_{\gamma, \delta}(H_2) = 1$ . Hence  $\hat{\mathcal{H}}_{\gamma, \delta} = \{H_1, H_2\}$ , and clearly  $\gamma, \delta \not\vdash_{RJ} H_1$ .  $\square$

### PROPOSITION A.3

Suppose  $\text{Fm}_\mathcal{D} \cap \text{Fm}_\mathcal{H} \neq \emptyset$ . Assume further that if  $T \models \neg\delta$  then  $r_\delta(H) = \infty$  for each  $H \in \mathcal{H}$ . Then  $\vdash_{RJ}$  satisfies the rules in Table 3.

PROOF. For (REF) we proceed as follows. First, note that from the assumptions on  $T$  it follows that, for every  $H \in \mathcal{H}$ , either  $T, H \models \delta$  or  $T, H \models \neg\delta$ . In the former case, we have  $r_\delta(H) = 0$  and  $H \in \hat{\mathcal{H}}_\delta$ , while in the latter  $r_\delta(H) > 0$ , hence  $H \notin \hat{\mathcal{H}}_\delta$ . In case that, at least for some  $H$ , we have  $T, H \models \delta$ , we then immediately obtain, for each  $H \in \hat{\mathcal{H}}_\delta$  that  $T, H \models \delta$  holds.

Assume now instead that, for all  $H \in \mathcal{H}$ , we have  $T, H \models \neg\delta$ . Since the hypotheses are exhaustive and mutually exclusive, this entails that  $T \models \neg\delta$ , hence  $r_\delta(H) = \infty$  for each  $H \in \mathcal{H}$  and by Lemma 2.4, we have that  $\delta \vdash_{RJ} \delta$ .

For (CUT), let us assume that (a)  $\Delta, \delta \vdash_{RJ} \psi$  and (b)  $\Delta \vdash_{RJ} \delta$ . If  $\hat{\mathcal{H}}_\Delta = \emptyset$  then by Lemma 2.4,  $\Delta \vdash_{RJ} \psi$  holds. So assume  $\hat{\mathcal{H}}_\Delta \neq \emptyset$  and let  $H \in \hat{\mathcal{H}}_\Delta$ . We need to show that  $T, H \models \psi$ . From (b) it follows that  $T, H \models \delta$  hence,  $r_\delta(H) = 0$ . But the latter means that  $r_{\Delta, \delta}(H) = r_\Delta(H) + r_\delta(H) = r_\Delta(H)$ . Hence, if  $H \in \hat{\mathcal{H}}_\Delta$ , we also have  $H \in \hat{\mathcal{H}}_{\Delta, \delta}$ . But by (a), this means that  $H \models \psi$ , thus showing our claim.

For (RMO), assume that  $\Delta \vdash_{RJ} \psi$  and  $\Delta \not\vdash_{RJ} \neg\varphi$ . Let us denote by  $H_\varphi$  all and only the hypothesis in  $\mathcal{H}$  that entail  $\varphi$ . Since the hypotheses in  $\mathcal{H}$  are mutually exclusive and exhaustive we have that the

set of all and only the hypotheses entailing  $\neg\varphi$  equals  $\mathcal{H} \setminus H_\varphi$ . Hence, from  $\Delta \vdash_{RJ} \neg\varphi$  it follows that  $\hat{\mathcal{H}}_\Delta \not\subseteq \mathcal{H} \setminus H_\varphi$ , i.e.  $\hat{\mathcal{H}}_\Delta \cap H_\varphi \neq \emptyset$ . Now, we claim that  $\hat{\mathcal{H}}_{\Delta,\varphi}$  is a subset of  $\hat{\mathcal{H}}_\Delta$ . From this, together with the premise  $\Delta \vdash_{RJ} \psi$ , it follows that, for every  $H \in \hat{\mathcal{H}}_{\Delta,\varphi}$  we have  $H \models \psi$ . But this amounts at saying that  $\Delta, \varphi \vdash_{RJ} \psi$ , thus showing that (RMO) holds.

Let us thus prove the claim that  $\hat{\mathcal{H}}_{\Delta,\varphi}$  is a subset of  $\hat{\mathcal{H}}_\Delta$ . Assume by contradiction that there is a  $H \in \hat{\mathcal{H}}_{\Delta,\varphi}$  such that  $H \notin \hat{\mathcal{H}}_\Delta$ . Now, recalling that  $\hat{\mathcal{H}}_\Delta \cap H_\varphi \neq \emptyset$ , pick  $H'$  in such a set. By the fact that  $H' \in H_\varphi$  and part (2) of Definition 2.1, we have  $r_\varphi(H') = 0$ , hence  $r_\varphi(H') \leq r_\varphi(H)$ . On the other hand since  $H' \in \hat{\mathcal{H}}_\Delta$  and  $H \notin \hat{\mathcal{H}}_\Delta$ , we have  $r_\Delta(H') < r_\Delta(H)$ . But these two facts entail that  $r_{\Delta,\varphi}(H') < r_{\Delta,\varphi}(H)$ , contradicting the fact that  $H \in \hat{\mathcal{H}}_{\Delta,\varphi}$ .

For (CMO) assume that  $\Delta \vdash_{RJ} \psi$  and  $\Delta \vdash_{RJ} \delta$ , i.e. that for any  $H \in \hat{\mathcal{H}}_\Delta$  we have both  $H \models \psi$  and  $H \models \delta$ . We now show that  $\hat{\mathcal{H}}_{\Delta,\delta} \subseteq \hat{\mathcal{H}}_\Delta$ .

Suppose that there is a  $H \in \hat{\mathcal{H}}_{\Delta,\delta}$  such that  $H \notin \hat{\mathcal{H}}_\Delta$ . Thus, there is a  $H' \in \hat{\mathcal{H}}_\Delta$  such that  $r_\Delta(H) > r_\Delta(H')$ . On the other hand, since  $H \in \hat{\mathcal{H}}_{\Delta,\delta}$  we have that  $r_{\Delta,\delta}(H) \leq r_{\Delta,\delta}(H')$ , from which it follows that  $r_\Delta(H) + r_\delta(H) \leq r_\Delta(H') + r_\delta(H')$  and  $r_\Delta(H) - r_\Delta(H') \leq r_\delta(H') - r_\delta(H)$ .

Since  $r_\Delta(H) > r_\Delta(H')$ , we have that  $0 < r_\Delta(H) - r_\Delta(H')$ . Consequently  $0 < r_\delta(H') - r_\delta(H)$  and  $r_\delta(H) < r_\delta(H')$ . Recall that  $\delta \in \text{Fm}_\mathcal{H}$  and in view of the assumption that the hypotheses are mutually exclusive and exhaustive, either  $T, H \models \neg\delta$  or  $T, H \models \delta$ . Hence we have that  $r_\delta(H) > 0$  iff  $H \models \neg\delta$ , otherwise  $r_\delta(H) = 0$ . Since  $H \models \delta$ , then  $H \not\models \neg\delta$  and  $r_\delta(H) = 0$ . Therefore, from  $r_\delta(H) < r_\delta(H')$  it follows that  $r_\delta(H') > 0$ . But  $r_\delta(H') > 0$  means in turn that  $H' \models \neg\delta$ , which contradicts the fact that  $H' \in \hat{\mathcal{H}}_\Delta$ , in view of the assumption  $\Delta \vdash_{RJ} \delta$ . We have thus shown that  $\hat{\mathcal{H}}_{\Delta,\delta} \subseteq \hat{\mathcal{H}}_\Delta$ . From this, the conclusion immediately follows. For, if  $\hat{\mathcal{H}}_{\Delta,\delta} = \emptyset$ , we immediately get  $\Delta, \delta \vdash_{RJ} \psi$ . Otherwise, for any  $H \in \hat{\mathcal{H}}_{\Delta,\delta}$ , we also have  $H \in \hat{\mathcal{H}}_\Delta$ , and by the assumption  $\Delta \vdash_{RJ} \psi$ , that  $H \models \psi$ . This finally shows that  $\Delta, \delta \vdash_{RJ} \psi$ .  $\square$

#### PROPOSITION A.4

Both  $(\text{AND})_I$  and  $(\text{AND})_I^{\text{con}}$  are invalid under  $\vdash_{\alpha RJ}$ .

PROOF. Let  $\mathcal{H} = \{H_1, H_2\}$  and  $\gamma, \delta \in \text{Fm}_\mathcal{D}$  such that  $\neg\gamma, \neg\delta \models \perp$ . Let  $\alpha_1 = \alpha_2 = 0.3$  and both  $t_{H_1}$  and  $t_{H_2}$  coincide with the identity function. Consider the following probability assignments:

$$\begin{aligned} P(\gamma \wedge \delta \mid H_1) &= 0.2 & P(\gamma \wedge \delta \mid H_2) &= 0.1 \\ P(\gamma \wedge \neg\delta \mid H_1) &= 0 & P(\gamma \wedge \neg\delta \mid H_2) &= 0.4 \\ P(\neg\gamma \wedge \delta \mid H_1) &= 0.8 & P(\neg\gamma \wedge \delta \mid H_2) &= 0.5 \end{aligned}$$

Since  $\neg\gamma, \neg\delta \models \perp$  we have

$$P(\neg\gamma \wedge \neg\delta \mid H_1) = P(\neg\gamma \wedge \neg\delta \mid H_2) = 0.$$

Now, by simple computations, we obtain

$$\begin{aligned} P(\gamma \mid H_1) &= 0.2 & P(\gamma \mid H_2) &= 0.5 \\ P(\delta \mid H_1) &= 1 & P(\delta \mid H_2) &= 0.6 \end{aligned}$$

### 30 A logical framework for data-driven reasoning

By the definition of  $r^\alpha$ , we thus get:

$$\begin{aligned} r_\gamma^\alpha(H_1) &= 0.8 & r_\gamma^\alpha(H_2) &= 0 \\ r_\delta^\alpha(H_1) &= 0 & r_\delta^\alpha(H_2) &= 0 \\ r_{\gamma \wedge \delta}^\alpha(H_1) &= 0.8 & r_{\gamma \wedge \delta}^\alpha(H_2) &= 0.9 \end{aligned}$$

We thus obtain

$$r_{\gamma, \delta}^\alpha(H_1) = r_\gamma^\alpha(H_1) + r_\delta^\alpha(H_1) = 0.8$$

and

$$r_{\gamma, \delta}^\alpha(H_2) = r_\gamma^\alpha(H_2) + r_\delta^\alpha(H_2) = 0$$

Thus  $\hat{\mathcal{H}}_{\gamma, \delta} = \{H_2\}$ , while on the other hand, since  $r_{\gamma \wedge \delta}^\alpha(H_1) < r_{\gamma \wedge \delta}^\alpha(H_2)$ , we also have  $\hat{\mathcal{H}}_{\gamma \wedge \delta} = \{H_1\}$ .

This shows that

$$\gamma, \delta \vdash_{\alpha RJ} H_2 \quad \gamma \wedge \delta \not\vdash_{\alpha RJ} H_2$$

thus proving that the  $(\text{AND})_I$  rule does not hold for  $\vdash_{\alpha RJ}$ .

The same example also shows that the converse  $(\text{AND})_I^{\text{con}}$  does not hold, since:

$$\gamma \wedge \delta \vdash_{\alpha RJ} H_1 \quad \gamma, \delta \not\vdash_{\alpha RJ} H_1 \quad \square$$

#### PROPOSITION A.5

$\vdash_{uRJ}$  satisfies the rules in Table 8.

PROOF. For (REF), (CUT), (CMO), (RMO) proceed as in Proposition 2.7.

For (UMO), assume that  $\Delta \vdash_{uRJ} \psi$  and  $\{\Delta \setminus \{\delta\} \vdash_{uRJ} \psi\}_{\delta \in \Delta}$ . By way of contradiction, let us assume  $\Delta, \varphi \not\vdash_{uRJ} \psi$ , i.e. that  $\hat{\mathcal{H}}_{\Delta, \varphi} \neq \emptyset$  (and thus  $r_{\Delta, \varphi}(H_{s^*}) < m + 1$ ) and in particular there is a  $H \in \hat{\mathcal{H}}_{\Delta, \varphi}$  such that  $H \not\models \psi$ . Note that  $H \notin \hat{\mathcal{H}}_\Delta$ , since otherwise by  $\Delta \vdash_{uRJ} \psi$ , we would immediately get  $H \models \psi$ .

Let us pick now any  $H' \in \hat{\mathcal{H}}_\Delta$ , so that  $r_\Delta(H) > r_\Delta(H')$  and  $H' \models \psi$  (since  $\Delta \vdash_{uRJ} \psi$ ). Note that, since  $H \in \hat{\mathcal{H}}_{\Delta, \varphi}$  and  $H' \in \hat{\mathcal{H}}_\Delta$ , both  $r_\Delta(H)$  and  $r_\Delta(H')$  are distinct from  $\infty$ . On the other hand, since  $H \in \hat{\mathcal{H}}_{\Delta, \varphi}$  we have  $r_{\Delta, \varphi}(H) \leq r_{\Delta, \varphi}(H')$ , hence  $r_\Delta(H) + r_\varphi(H) \leq r_\Delta(H') + r_\varphi(H')$  and  $r_\varphi(H) - r_\varphi(H') \leq r_\Delta(H') - r_\Delta(H) < 0$ . Thus in particular we obtain  $r_\varphi(H) < r_\varphi(H')$ . But this can only occur if  $r_\varphi(H) = 0$  and  $r_\varphi(H') = 1$ . And this means that  $r_{\Delta, \varphi}(H) = r_\Delta(H)$  while  $r_{\Delta, \varphi}(H') = r_\Delta(H') + 1$  and since  $r_{\Delta, \varphi}(H) \leq r_{\Delta, \varphi}(H')$ , we have  $r_\Delta(H) \leq r_\Delta(H') + 1$ . Hence we have both  $r_\Delta(H) > r_\Delta(H')$  and  $r_\Delta(H) \leq r_\Delta(H') + 1$ , i.e.  $r_\Delta(H) = r_\Delta(H') + 1$ .

This means that there is a  $\bar{\delta} \in \Delta$  rejecting  $H$  and not rejecting  $H'$ . Let us consider  $\Delta \setminus \{\bar{\delta}\}$ . By our choice of  $\bar{\delta}$ , we will now have  $r_{\Delta \setminus \{\bar{\delta}\}}(H') = r_{\Delta \setminus \{\bar{\delta}\}}(H)$ , hence  $H' \in \hat{\mathcal{H}}_{\Delta \setminus \{\bar{\delta}\}}$ . But this, by the assumption,  $\{\Delta \setminus \{\delta\} \vdash_{uRJ} \psi\}_{\delta \in \Delta}$ , entails  $H' \models \psi$ , which is the desired contradiction.  $\square$

#### LEMMA A.6

$\vdash_{uRJ}$  satisfies

$$\frac{\Delta, \gamma \vdash \psi \quad \Delta, \delta \vdash \psi}{\Delta, \gamma \vee \delta \vdash \psi.} \quad (\text{OR})$$

PROOF. Let  $\Delta, \gamma \vdash_{uRJ} \psi$  and  $\Delta, \delta \vdash_{uRJ} \psi$ . To show  $\Delta, \gamma \vee \delta \vdash_{uRJ} \psi$ , let  $H \in \hat{\mathcal{H}}_{\gamma \vee \delta}$ . If either  $H \in \hat{\mathcal{H}}_{\gamma}$  or  $H \in \hat{\mathcal{H}}_{\delta}$ , we immediately get, by the assumption, that  $H \models \psi$ . We thus assume  $H \notin \hat{\mathcal{H}}_{\gamma}$  and  $H \notin \hat{\mathcal{H}}_{\delta}$ , and derive a contradiction. By  $H \notin \hat{\mathcal{H}}_{\delta}$ , we obtain that there exists a  $H'$  such that

$$r_{\Delta, \gamma}(H') = r_{\Delta}(H') + r_{\gamma}(H') < r_{\Delta}(H) + r_{\gamma}(H) = r_{\Delta, \gamma}(H)$$

and by  $H \notin \hat{\mathcal{H}}_{\psi}$ , analogously, we obtain that there exists a  $H''$  such that

$$r_{\Delta, \delta}(H'') = r_{\Delta}(H'') + r_{\delta}(H'') < r_{\Delta}(H) + r_{\delta}(H) = r_{\Delta, \delta}(H)$$

However, since  $H \in \hat{\mathcal{H}}_{\gamma \vee \delta}$  we have:

$$r_{\Delta, \gamma \vee \delta}(H) = r_{\Delta}(H) + r_{\gamma \vee \delta}(H) \leq r_{\Delta}(H') + r_{\gamma \vee \delta}(H') = r_{\Delta, \gamma \vee \delta}(H')$$

and

$$r_{\Delta, \gamma \vee \delta}(H) = r_{\Delta}(H) + r_{\gamma \vee \delta}(H) \leq r_{\Delta}(H'') + r_{\gamma \vee \delta}(H'') = r_{\Delta, \gamma \vee \delta}(H'').$$

Now, note that  $r_{\gamma \vee \delta}(H) = 1$  if and only if  $\gamma \vee \delta \models \neg H$ . Hence  $r_{\gamma \vee \delta}(H) = 1$  if  $\gamma \models \neg H$  and  $\delta \models \neg H$ , and it is 0 otherwise. In other words,  $r_{\gamma \vee \delta}(H) = \min(r_{\gamma}(H), r_{\delta}(H))$ , and the same clearly holds for  $H'$  and  $H''$ . Thus, combining our previous inequalities, we obtain

$$r_{\Delta}(H) + r_{\gamma \vee \delta}(H) \leq r_{\Delta}(H') + r_{\gamma \vee \delta}(H') \leq r_{\Delta}(H') + r_{\gamma}(H') < r_{\Delta}(H) + r_{\gamma}(H)$$

and similarly

$$r_{\Delta}(H) + r_{\gamma \vee \delta}(H) \leq r_{\Delta}(H'') + r_{\gamma \vee \delta}(H'') \leq r_{\Delta}(H'') + r_{\delta}(H'') < r_{\Delta}(H) + r_{\delta}(H)$$

But now, since  $r_{\gamma \vee \delta}(H) = \min(r_{\gamma}(H), r_{\delta}(H))$  it is either equal to  $r_{\gamma}(H)$  or to  $r_{\delta}(H)$ , contradiction.  $\square$

LEMMA A.7

$\vdash_{uRJ}$  satisfies

$$\frac{\Delta, \varphi \vdash_{uRJ} \psi \quad \Delta \vdash_{uRJ} \neg \psi}{\Delta \vdash_{uRJ} \neg \varphi.} \text{ (MT)}$$

PROOF. Assume the two premisses of MT:

(A1)  $\Delta, \varphi \vdash_{uRJ} \psi$ , i.e. for each  $H \in \hat{\mathcal{H}}_{\Delta, \varphi}$ ,  $H \models \psi$ , and that

(A2)  $\Delta \vdash_{uRJ} \neg \psi$ , i.e. for each  $H \in \hat{\mathcal{H}}_{\Delta}$ ,  $H \models \neg \psi$ .

We show that for each  $H \in \hat{\mathcal{H}}_{\Delta}$ ,  $H \models \neg \varphi$ . Suppose, on the contrary, that there exists  $H' \in \hat{\mathcal{H}}_{\Delta}$  such that  $H' \not\models \neg \varphi$ , i.e.  $\varphi$  does not reject  $H'$ . Therefore  $H' \in \hat{\mathcal{H}}_{\Delta, \varphi}$  and  $H' \models \psi$ , but this is in contradiction with (A2).  $\square$

Suppose  $\Delta \vdash_{uRJ} \psi$  and  $r_{\Delta}(H)$  is finite for each  $H \in \mathcal{H}$ . Then there is a derivation of it using the rules in Table 7 and Table 8.

PROOF. We will show how to find a derivation of  $\Delta' \vdash_{uRJ} \neg H_j$ , proceeding by induction on the number of formulas occurring in  $\Delta'$ . Henceforth, in the applications of (UMO) and (RMO) we will not explicitly consider the additional condition that  $r_{\Delta, \varphi}(H)$  is finite for each  $H \in \mathcal{H}$  since this will always hold, under our assumption that  $r'_{\Delta}(H)$  is finite for each  $H$ .

For the base case, when  $\Delta'$  is composed of only one formula, it has to be of the form  $\neg H_j \vdash_{uRJ} \neg H_j$ , which is clearly an instance of (REF). That this can be the only case when the size of  $\Delta'$  is one,

follows from the fact that, for any  $i \neq j$  we would have that  $H_j \in \hat{\mathcal{H}}_{\neg H_i}$  and  $H_j \not\models \neg H_j$ . Hence  $\neg H_i \not\sim_{uRJ} \neg H_j$  if  $i \neq j$ .

Now let us consider the inductive step. By Definition 4.4 and the normal form of the formulas in  $\Delta'$  we have two possibilities:

- (1)  $\hat{\mathcal{H}}_{\Delta'} = \{H \mid \neg H \notin \Delta'\}$  i.e. the least rejected hypotheses are those that do not occur in  $\Delta'$  at all, or
- (2)  $\hat{\mathcal{H}}_{\Delta'} = \{H_k \mid l_k \neq 0 \text{ and } l_k \text{ is a minimal non-zero index among the } l_1, \dots, l_n\}$ .

In case (1), for any  $H \in \hat{\mathcal{H}}_{\Delta'}$ , we have  $\neg H \notin \Delta'$ , i.e.  $r_{\Delta'}(H) = 0$ . However, we have that  $H_j \notin \hat{\mathcal{H}}_{\Delta'}$ . Indeed, if we had  $H_j \in \hat{\mathcal{H}}_{\Delta'}$ , from  $\Delta' \vdash_{uRJ} \neg H_j$ , we would get that  $H_j \models \neg H_j$ , which is absurd.

Now, if there are no hypotheses in  $\Delta'$  distinct from  $H_j$ , this means that our consequence is of the form  $(\neg H_j)^{l_j} \vdash_{uRJ} \neg H_j$ . Hence we could reduce its size by the following application of RMO.

$$\frac{\neg H_j^{l_j-1} \vdash_{uRJ} \neg H_j \quad (\neg H_j)^{l_j-1} \not\sim_{uRJ} \neg \neg H_j}{(\neg H_j)^{l_j} \vdash_{uRJ} \neg H_j} \text{ (RMO)}$$

Now upon removal of any formula  $\neg H'$  from  $\Delta'$  distinct from  $\neg H_j$ , we will still have  $r_{\Delta' \setminus \{\neg H'\}}(H) = 0$ , for any  $H \in \hat{\mathcal{H}}_{\Delta'}$ . This means that, if  $H \in \hat{\mathcal{H}}_{\Delta'}$ , then still  $H \in \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H'\}}$ . Recall now that  $T_{\Delta}, H \models \neg H'$  for any  $\neg H' \in \Delta'$ , and  $H \models \neg H_j$ .

On the other hand, by assumption,  $\hat{\mathcal{H}}_{\Delta'}$  does not contain  $H_j$ , and since  $H' \neq H_j$ , we may safely assume that  $H_j \notin \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H'\}}$ . This means that still  $\Delta' \setminus \{\neg H'\} \vdash_{uRJ} H_j$ .

Hence, we have shown that both  $\Delta' \setminus \{\neg H'\} \not\vdash_{uRJ} H'$  and  $\Delta' \setminus \{\neg H'\} \vdash_{uRJ} H_j$ . We then get  $\Delta' \vdash_{uRJ} \neg H_j$  by applying (RMO) as follows:

$$\frac{\Delta' \setminus \{\neg H'\} \vdash_{uRJ} \neg H_j \quad \Delta' \setminus \{\neg H'\} \not\sim_{uRJ} \neg \neg H'}{\Delta' \vdash_{uRJ} \neg H_j} \text{ (RMO)}.$$

Let us now consider case (2), where the least rejected hypotheses occur (negated) in  $\Delta'$ , i.e. for any  $H \in \hat{\mathcal{H}}_{\Delta'}$ , we have  $\neg H \in \Delta'$ , hence  $r_{\Delta'}(H) \neq 0$ . Since  $H_j \not\models \neg H_j$  and  $\Delta' \vdash_{uRJ} \neg H_j$  we may again exclude that  $H_j$  is any of the least rejected hypotheses.

Let us first consider the subcase (2a) where  $H_j$  is the maximally rejected hypothesis. We need to distinguish three sub(sub)cases.

(2a') If in  $\Delta'$  there is another maximally rejected hypothesis, say  $H_p$ , in addition to  $H_j$ , we use (RMO) to remove (reasoning backwards) one of its occurrences.

First, note that since  $\Delta' \vdash_{uRJ} \neg H_j$ , we have, for any  $H \in \hat{\mathcal{H}}_{\Delta'}$ , that  $H \models \neg H_j$ , and in particular  $H_j \notin \hat{\mathcal{H}}_{\Delta'}$ . The removal of an occurrence of a maximally rejected  $\neg H_p$  will not affect this. In other words, we still have  $H_j \notin \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_p\}}$  hence  $\Delta' \setminus \{\neg H_p\} \vdash_{uRJ} \neg H_j$ . On the other hand, recall that by assumption of case (2) the least rejected hypotheses, say  $H_k$ , occur in  $\Delta'$ . Upon removal of an occurrence of the maximally rejected hypothesis  $H_p$ , we still have  $H_k \in \hat{\mathcal{H}}_{\Delta' \setminus \{\neg H_p\}}$ . Since  $H_k \models \neg H_p$ , we will then have  $\Delta' \setminus \{\neg H_p\} \not\sim_{uRJ} H_p$ . Hence we may derive  $\Delta' \vdash_{uRJ} \neg H_j$  as follows:

$$\frac{\Delta' \setminus \{\neg H_p\} \vdash_{uRJ} \neg H_j \quad \Delta' \setminus \{\neg H_p\} \not\sim_{uRJ} \neg \neg H_p}{\Delta' \vdash_{uRJ} \neg H_j} \text{ (RMO)}.$$

(2a'') Assume now that: (i)  $H_j$  is a maximally rejected hypothesis (as per assumption of case (2a)); (ii)  $H_j$  is the unique maximally rejected hypothesis (in contrast to case (2a')); (iii)  $H_k$  is a least rejected hypothesis, that by assumption (2) occurs (negated) in  $\Delta'$ .

This means that we have  $r_{\Delta'}(H_j) > r_{\Delta'}(H_k)$ . Note that  $r_{\Delta'}(H_j) = r_{\Delta' \setminus \{-H_k\}}(H_j)$  and  $r_{\Delta'}(H_k) = r_{\Delta' \setminus \{-H_k\}}(H_k) + 1$ , hence

$$r_{\Delta' \setminus \{-H_k\}}(H_j) > r_{\Delta' \setminus \{-H_k\}}(H_k) + 1 > r_{\Delta' \setminus \{-H_k\}}(H_k). \quad (*)$$

This entails  $H_j \notin \hat{\mathcal{H}}_{\Delta' \setminus \{-H_k\}}$ , hence  $\Delta' \setminus \{-H_k\} \vdash_{uRJ} \neg H_j$ . However, for any  $\neg H \in \Delta' \setminus \{H_k\}$ , we have either

$$r_{\Delta' \setminus \{-H_k, \neg H\}}(H_j) = r_{\Delta' \setminus \{-H_k\}}(H_j) - 1$$

or

$$r_{\Delta' \setminus \{-H_k, \neg H\}}(H_j) = r_{\Delta' \setminus \{-H_k\}}(H_j)$$

depending on whether  $\neg H \models \neg H_j$  (i.e.  $H$  is actually  $H_j$ ) or not. In both cases we have

$$r_{\Delta' \setminus \{-H_k, \neg H\}}(H_j) \geq r_{\Delta' \setminus \{-H_k\}}(H_j) - 1.$$

Hence we obtain

$$r_{\Delta' \setminus \{-H_k, \neg H\}}(H_j) \geq r_{\Delta' \setminus \{-H_k\}}(H_j) - 1 > r_{\Delta' \setminus \{-H_k\}}(H_k) \geq r_{\Delta' \setminus \{-H_k, \neg H\}}(H_k)$$

where the strict inequality follows from (\*), and the last inequality follows directly from the definition of the degree of rejection  $r$ .

This shows that  $H_j \notin \hat{\mathcal{H}}_{\Delta' \setminus \{-H_k, \neg H\}}$  for any choice of  $\neg H$  (including  $\neg H_j$ ). Hence, for any  $\neg H \in \Delta'$ :

$$\Delta' \setminus \{-H_k, \neg H\} \vdash_{uRJ} \neg H_j.$$

We can thus derive  $\Delta' \vdash_{uRJ} \neg H_j$  using the rule application:

$$\frac{\Delta' \setminus \{-H_k\} \vdash_{uRJ} \neg H_j \quad \{\Delta' \setminus \{-H_k, \neg H\} \vdash_{uRJ} \neg H_j\}_{\neg H \in \Delta'}}{\Delta' \vdash_{uRJ} \neg H_j} \text{ (UMO)}$$

and the above reasoning, showing that its premises hold.

Finally, our last subcase of (2a) is (2a''), when there is no other hypotheses but  $\neg H_j$  in  $\Delta'$ , i.e.  $\Delta' \vdash_{uRJ} \neg H_j$  is actually of the form  $\neg H_j^j \vdash_{uRJ} \neg H_j$ . We derive it, starting from  $\neg H_j \vdash_{uRJ} \neg H_j$ , by repeated applications of (RMO), beginning with

$$\frac{\neg H_j \vdash_{uRJ} \neg H_j \quad \neg H_j \not\vdash_{uRJ} \neg \neg H_j}{\neg H_j, \neg H_j \vdash_{uRJ} \neg H_j} \text{ (RMO)}.$$

Consider now case (2b), where  $H_j$  is not a maximally rejected hypotheses, and let  $H_p$  be any maximally rejected hypothesis. Since  $\Delta' \vdash_{uRJ} \neg H_j$ , we know that  $H_j \notin \hat{\mathcal{H}}_{\Delta'}$ , i.e.  $H_j$  is not a least rejected hypothesis either. Upon removal of one occurrence of  $\neg H_p$ ,  $H_j$  will still not be a least rejected hypothesis, hence  $H_j \notin \hat{\mathcal{H}}_{\Delta' \setminus \{-H_p\}}$  and  $\Delta' \setminus \{-H_p\} \vdash_{uRJ} \neg H_j$ . Moreover, there will be at least a least rejected hypothesis in  $\hat{\mathcal{H}}_{\Delta' \setminus \{-H_p\}}$  that is distinct from  $H_p$ , hence  $\Delta' \setminus \{-H_p\} \not\vdash_{uRJ} H_p$ .

This means that we may consider the following application of (RMO):

$$\frac{\Delta' \setminus \{-H_p\} \vdash_{uRJ} \neg H_j \quad \Delta' \setminus \{-H_p\} \not\vdash_{uRJ} \neg \neg H_p}{\Delta' \vdash_{uRJ} \neg H_j} \text{ (RMO)},$$

completing our proof.  $\square$

## References

- [1] Adams E. *A Primer of Probability Logic*. CSLI Publications, 1998.
- [2] Baldi P, Corsi EA, Hosni H. Logical desiderata on statistical inference. In Antunes H, Freire A, Rodrigues A (eds.), *Forthcoming in "Walter Carnielli on Reasoning, Paraconsistency, and Probability"*, Springer's series "Outstanding Contributions to Logic". Springer, 2023, 1–24.
- [3] Baltag A, Rafiee Rad S, Smets S. Tracking probabilistic truths: a logic for statistical learning. *Synthese* 2021; **199**:9041–87. <https://doi.org/10.1007/s11229-021-03193-6>.
- [4] Bookstein FL. *Measuring and Reasoning*. Cambridge University Press, 2014.
- [5] Borges W, Stern, JM. The rules of logic composition for the Bayesian epistemic e-values. *Log J IGPL* 2007; **15**:401–20. <https://doi.org/10.1093/jigpal/jzm032>.
- [6] Bosley R. Modus Tollens. *Notre Dame J Form Log* 1979; **20**:103–11. <https://doi.org/10.1305/ndjfl/1093882408>.
- [7] Carnielli W, Coniglio M, D'Ottaviano IML. *Paraconsistency: The Logical Way to the Inconsistent*. Marcel Dekker, 2002.
- [8] Carnielli W, Coniglio ME, Marcos J. 2007. Logics of formal inconsistency. In Gabbay and Guenther (eds), *Handbook of Philosophical Logic*, vol. 14. Springer.
- [9] Carnielli W, Coniglio M. *Paraconsistent Logic: Consistency, Contradiction and Negation*. Springer, 2016.
- [10] Carnielli W, Rodrigues A. An epistemic approach to paraconsistency: a logic of evidence and truth. *Synthese* 2017; **196**:3789–813.
- [11] Rodrigues A, Bueno-Soler A, Carnielli W. Measuring evidence: a probabilistic approach to an extension of Belnap–Dunn logic. *Synthese* 2021; **198**:5451–80.
- [12] Chow SL. *Statistical Significance*. Sage, 1996.
- [13] Dickson M, Baird D. Significance testing. In *Handbook of the Philosophy of Science, vol. 7: Philosophy of Statistics*, vol. 7. Elsevier, 2011, 199–229.
- [14] Esteves L, Izbicki R, Stern J. *et al.* The logical consistency of simultaneous agnostic hypothesis tests. *Entropy* 2016; **18**:1–22. <https://doi.org/10.3390/e18070256>.
- [15] Falk R, Greenbaum CW. Significance tests die hard: the amazing persistence of a probabilistic misconception. *Theory Psychol* 1995; **5**:75–98. <https://doi.org/10.1177/0959354395051004>.
- [16] Fisher RA. Statistical tests. *Nature* 1935; **136**:474. <https://doi.org/10.1038/136474b0>.
- [17] Fisher RA. *Statistical Methods for Research Workers*, 12th edn. New York: Hafner Publishing Company, 1954.
- [18] Fisher RA. *Statistical Methods and Statistical Inference*. Oliver and Boyd, 1956.
- [19] Gabbay DM. *Logic for Artificial Intelligence & Information Technology*. London: College Publications, 2007.
- [20] Galatos N, Jipsen P, Kowalski T. *et al.* *Residuated Lattices: An Algebraic Glimpse at Substructural Logics*. Elsevier, 2007.
- [21] Giles G. Semantics for fuzzy reasoning. *Int J Man-Machine Studies* 1982; **17**:401–15. [https://doi.org/10.1016/S0020-7373\(82\)80041-2](https://doi.org/10.1016/S0020-7373(82)80041-2).
- [22] Halpern JY. *Reasoning About Uncertainty*. MIT Press, 2003.
- [23] Hawthorne J, Makinson D. The quantitative/qualitative watershed for rules of uncertain inference. *Stud Logica* 2007; **86**:247–97. <https://doi.org/10.1007/s11225-007-9061-x>.
- [24] Hey T, Butler K, Jackson S. *et al.* Machine learning and big scientific data. *Philos Trans R Soc A Math Phys Eng Sci* 2020; **378**:20190054. <https://doi.org/10.1098/rsta.2019.0054>.
- [25] Hosni H, Landes J. Logical perspectives on the foundations of probability. *Open Math* 2023; **21**:20220598. <https://doi.org/10.1515/math-2022-0598>.

- [26] Howson C, Urbach P. *Scientific Reasoning: The Bayesian Approach*. Open Court, 1993.
- [27] Izbicki R, Esteves L. Logical consistency in simultaneous statistical test procedures. *Log J IGPL* 2015; **23**:732–58. <https://doi.org/10.1093/jigpal/jzv027>.
- [28] Kinraide TB, Denison R. Strong inference: the way of science. *Am Biol Teach* 2003; **65**:419–24. <https://doi.org/10.2307/4451529>.
- [29] Klein D, Majer O, Rafiee Rad S. Probabilities with gaps and gluts. *J Philos Log* 2021; **50**:1107–41. <https://doi.org/10.1007/s10992-021-09592-x>.
- [30] Kraus S, Lehmann D, Magidor M. Nonmonotonic reasoning, preferential models and cumulative logics. *Artif Intell* 1990; **44**:167–207. [https://doi.org/10.1016/0004-3702\(90\)90101-5](https://doi.org/10.1016/0004-3702(90)90101-5).
- [31] Krueger J. Null hypothesis significance testing: on the survival of a flawed method. *Am Psychol* 2001; **56**:16–26. <https://doi.org/10.1037/0003-066X.56.1.16>.
- [32] Kyburg HE, Man Teng C. *Uncertain Inference*. Cambridge University Press, 2001.
- [33] Kyburg HE, Man Teng C. Nonmonotonic logic and statistical inference. *Comput Intell* 2006; **22**:26–51. <https://doi.org/10.1111/j.1467-8640.2006.00272.x>.
- [34] Lehmann D, Magidor M. What does a conditional knowledge base entail? 1992; **55**:1–60. [https://doi.org/10.1016/0004-3702\(92\)90041-U](https://doi.org/10.1016/0004-3702(92)90041-U).
- [35] Lehmann E. L. *Fisher, Neyman, and the Creation of Classical Statistics*. 2011.
- [36] Makinson D. *Bridges From Classical to Non-Monotonic Logic*. London: College Publications, 2005.
- [37] Makinson D. Logical questions behind the lottery and preface paradoxes: lossy rules for uncertain inference. *Synthese* 2012; **186**:511–29. <https://doi.org/10.1007/s11229-011-9997-2>.
- [38] Marquis P, Papini O, Prade H. *A Guided Tour of Artificial Intelligence Research 1: Knowledge Representation, Reasoning and Learning*. Springer, 2020.
- [39] Mayo DG. *Statistical Inference as Severe Testing: How to Get Beyond the Statistics Wars*. Cambridge University Press, 2018.
- [40] Mierzewski K. Probabilistic stability: dynamics, non-monotonic logics, and stable revision. MSc Thesis, Amsterdam University, 2018.
- [41] Mundici D. The logic of Ulam’s game with lies. In Bicchieri C, Dalla Chiara ML (eds.), *Knowledge, Belief, and Strategic Interaction*. Cambridge Studies in Probability, Induction and Decision Theory. Cambridge University Press, 1992, 275–84. <https://doi.org/10.1017/CBO9780511983474.017>.
- [42] Mundici D. Ulam games, Lukasiewicz logic, and AF C\*-algebras. *Fundam Inform* 1993; **18**:151–61. <https://doi.org/10.3233/FI-1993-182-405>.
- [43] Omori H, Wansing H. 40 years of FDE: an introductory overview. *Stud Logica* 2017; **105**:1021–49. <https://doi.org/10.1007/s11225-017-9748-6>.
- [44] Pearl J. *Probabilistic Reasoning in Intelligent Systems: Networks of Plausible Inference*. Morgan Kaufmann, 1988.
- [45] Pereira CAB, Stern JM. The e-value: a fully Bayesian significance measure for precise statistical hypotheses and its research program. *Sao Paulo J Math Sci* 2022; **16**:566–84. <https://doi.org/10.1007/s40863-020-00171-7>.
- [46] Pollard P, Richardson JT. On the probability of making type I errors. *Psychol Bull* 1987; **102**:159–63.
- [47] Platt JR. Strong inference. *Science* 1964; **146**:347–53. <https://doi.org/10.1126/science.146.3642.347>.
- [48] Rényi A. On a problem of information theory. *MTA Mat Kut Int Kozl B* 1961; **6**:505–16.
- [49] Reiter R, Criscuolo G. On interacting defaults. *IJCAI Int Joint Conf Artif Intell* 1981;413–20.

- [50] Ritchie S. *Science Fictions. How Fraud, Bias, Negligence and Hype Undermine the Search for Truth*. Macmillan, 2020.
- [51] Royall R. *Statistical Evidence: A Likelihood Paradigm*. Chapman & Hall, 1997.
- [52] Shoham Y. Nonmonotonic logics: meaning and utility. *Proc IJCAI-87* 1987;388–92.
- [53] Stern J. M. Paraconsistent sensitivity analysis for Bayesian significance tests. *Brazil Symp Artif Intell* 2004;134–43. [https://doi.org/10.1007/978-3-540-28645-5\\_14](https://doi.org/10.1007/978-3-540-28645-5_14).
- [54] Stern JM, Izbicki R, Esteves LG. *et al.* Logically-consistent hypothesis testing and the hexagon of oppositions. *Log J IGPL* 2017; **25**:741–57. <https://doi.org/10.1093/jigpal/jzx024>.
- [55] Szucs D, Ioannidis JPA. When null hypothesis significance testing is unsuitable for research: a reassessment. *Front Hum Neurosci* 2017; **11**:390. <https://doi.org/10.3389/fnhum.2017.00390>.
- [56] Ulam SM. *Adventures of a Mathematician*. Univ of California Press, 1991.
- [57] Wagner CG. Modus Tollens probabilized. *Brit J Philos Sci* 2004; **55**:747–53. <https://doi.org/10.1093/bjps/55.4.747>.
- [58] Wasserstein RL, Lazar NA. The ASA’s statement on p-values: context, process, and purpose. *Am Stat* 2016; **70**:129–33. <https://doi.org/10.1080/00031305.2016.1154108>.
- [59] Wasserstein RL, Schirm AL, Lazar NA. Moving to a world beyond “p < 0.05”. *Am Stat* 2019; **73**:1–19. <https://doi.org/10.1080/00031305.2019.1583913>.

Received 6 May 2024