

COMENIUS UNIVERSITY IN BRATISLAVA
FACULTY OF MATHEMATICS, PHYSICS AND INFORMATICS

COHERENCE-BASED BELIEF REVISION IN
EVOLUTIONARY COMPUTING: EVALUATING
EPISTEMIC RULES BY SPEED AND ACCURACY
MASTER'S THESIS

2026

BC. DÁVID BARTOŠ

COMENIUS UNIVERSITY IN BRATISLAVA
FACULTY OF MATHEMATICS, PHYSICS AND INFORMATICS

COHERENCE-BASED BELIEF REVISION IN
EVOLUTIONARY COMPUTING: EVALUATING
EPISTEMIC RULES BY SPEED AND ACCURACY
MASTER'S THESIS

Study Programme: Cognitive Science
Field of Study: Computer Science
Department: Department of Applied Informatics
Supervisor: Asst. Prof. Dr. Borut Trpin
Consultant: prof. Ing. Igor Farkaš, Dr.

Bratislava, 2026
Bc. Dávid Bartoš



ZADANIE ZÁVEREČNEJ PRÁCE

Meno a priezvisko študenta: Bc. Dávid Bartoš
Študijný program: kognitívna veda (Jednoodborové štúdium, magisterský II. st., denná forma)
Študijný odbor: informatika
Typ záverečnej práce: diplomová
Jazyk záverečnej práce: anglický
Sekundárny jazyk: slovenský

Názov: Coherence-Based Belief Revision in Evolutionary Computing: Evaluating Epistemic Rules by Speed and Accuracy
Revízia presvedčení založená na koherencii v evolučných výpočtoch: Hodnotenie epistemických pravidiel rýchlosťou a presnosťou

Anotácia: Práca skúma revíziu presvedčení v podmienkach neistoty porovnaním niekoľkých aktualizáčnych pravidiel vo výpočtovej simulácii. Práca kombinuje filozofiu vedy, formálnu epistemológiu, teóriu pravdepodobnosti, kognitívnu vedu a evolučné výpočty s cieľom študovať, ako rôzne epistemologické stratégie sledujú pravdu v podmienkach neistoty.

Cieľ: 1. Porovnajte rôzne pravidlá revízie presvedčení za podmienok neistoty v riadenej výpočtovej simulácii.
2. Preskúmajte, ako fungujú Bayesovská aktualizácia, vysvetľujúca aktualizácia, aktualizácia založená na koherencii a evolučný výber, keď agenti opakovane aktualizujú svoje presvedčenia na základe zašumenej evidencie.

Literatúra: Douven, I. (2019). Optimizing group learning: An evolutionary computing approach. *Artificial Intelligence*, 275, 235–251. <https://doi.org/10.1016/j.artint.2019.06.002>
Douven, I., & Schupbach, J. N. (2015). Probabilistic alternatives to Bayesianism: The case of explanationism. *Frontiers in Psychology*, 6, <https://doi.org/10.3389/fpsyg.2015.00459>
Shogenji, T. (1999). Is coherence truth conducive? *Analysis*, 59(4), 338–345. <https://doi.org/10.1093/analys/59.4.338>

Vedúci: Borut Trpin, PhD.
Konzultant: prof. Ing. Igor Farkaš, Dr.
Katedra: FMFI.KAI - Katedra aplikovanej informatiky
Vedúci katedry: doc. RNDr. Tatiana Jajcayová, PhD.
Dátum zadania: 13.05.2024

Dátum schválenia: 20.05.2024

prof. Ing. Igor Farkaš, Dr.
garant študijného programu



THESIS ASSIGNMENT

Name and Surname: Bc. Dávid Bartoš
Study programme: Cognitive Science (Single degree study, master II. deg., full time form)
Field of Study: Computer Science
Type of Thesis: Diploma Thesis
Language of Thesis: English
Secondary language: Slovak

Title: Coherence-Based Belief Revision in Evolutionary Computing: Evaluating Epistemic Rules by Speed and Accuracy

Annotation: The thesis investigates belief revision under uncertainty by comparing several update rules in a computational simulation. The thesis combines philosophy of science, formal epistemology, probability theory, cognitive science, and evolutionary computation to study how different epistemic strategies track truth under uncertainty.

Aim:

1. Compare different rules of belief revision under uncertainty in a controlled computational simulation.
2. Examine how Bayesian updating, explanatory updating, coherence-based updating, and evolutionary selection perform when agents repeatedly update their beliefs on the basis of noisy evidence.

Literature: Douven, I. (2019). Optimizing group learning: An evolutionary computing approach. *Artificial Intelligence*, 275, 235–251. <https://doi.org/10.1016/j.artint.2019.06.002>
Douven, I., & Schupbach, J. N. (2015). Probabilistic alternatives to Bayesianism: The case of explanationism. *Frontiers in Psychology*, 6, <https://doi.org/10.3389/fpsyg.2015.00459>
Shogenji, T. (1999). Is coherence truth conducive? *Analysis*, 59(4), 338–345. <https://doi.org/10.1093/analys/59.4.338>

Supervisor: Borut Trpin, PhD.
Consultant: prof. Ing. Igor Farkaš, Dr.
Department: FMFI.KAI - Department of Applied Informatics
Head of department: doc. RNDr. Tatiana Jajcayová, PhD.

Assigned: 13.05.2024

Approved: 20.05.2024
prof. Ing. Igor Farkaš, Dr.
Guarantor of Study Programme

Student

Supervisor

Acknowledgments: Here I would like to express my sincere gratitude to my supervisor, Asst. Prof. Dr. Borut Trpin, for the considerable time he devoted to online consultations throughout the preparation of this thesis. His patient explanations of the mathematical foundations and computational implementation of the study were invaluable. I am also deeply grateful for his detailed feedback on the content of the thesis, which helped me clarify and improve both the theoretical discussion and the presentation of the computational work.

I would also like to thank my consultant, prof. Ing. Igor Farkaš, Dr., for creating the LaTeX template used for this thesis, for his assistance with its preparation and formatting, and for his helpful suggestions and textual revisions during our online working sessions.

Abstract

This thesis compares seven belief-updating methods in a computer simulation designed to test how quickly and accurately they identify the correct hypothesis from repeated evidence. Each method was represented by agents that estimated the coin's true bias, and better-performing agents were more likely to remain in the population over successive generations. Performance was assessed by combining two demands: reaching strong confidence in the correct answer quickly and maintaining accurate probability estimates throughout the process. The results revealed two stable patterns. A method that gave extra support to the hypothesis closest to the observed frequency dominated in seven of the ten tested conditions. Methods that rewarded coherence, understood as how well the accumulated evidence fits together, dominated in the three intermediate conditions. Changes in the amount of evidence, the confidence requirement, and the strength of selection mainly affected how quickly unsuccessful methods disappeared rather than which methods ultimately survived. Standard probability updating and two methods based on explanatory strength did not survive to the final generation in any tested condition, although one method based on coherence often remained competitive for longer. The findings therefore do not identify a universally best method. Instead, they show that success depends on the structure of the evidence and that explanation and coherence can provide useful additions to standard probability-based belief revision in some environments. The study also demonstrates how philosophical ideas about rational reasoning can be translated into testable computer models.

Keywords: belief revision, uncertainty, coherence, evolutionary computing, probabilistic reasoning

Abstrakt

Táto diplomová práca porovnáva sedem metód aktualizácie presvedčení v počítačovej simulácii navrhutej na overenie toho, ako rýchlo a presne dokážu na základe opakovane získavaných dôkazov určiť správnu hypotézu. Každú metódu reprezentovali agenti, ktorí odhadovali skutočnú mieru skreslenia mince, pričom výkonnejší agenti mali väčšiu pravdepodobnosť zotrvať v populácii počas nasledujúcich generácií. Výkon sa hodnotil spojením dvoch požiadaviek: rýchleho dosiahnutia vysokej miery istoty v správnu odpoveď a zachovania presných odhadov pravdepodobnosti počas celého procesu. Výsledky odhalili dva stabilné vzorce. Metóda, ktorá poskytovala dodatočnú podporu hypotéze najbližšej k pozorovanej frekvencii, dominovala v siedmich z desiatich skúmaných podmienok. Metódy zvýhodňujúce koherenciu, chápanú ako miera vzájomného súladu nahromadených dôkazov, dominovali v troch stredných podmienkach. Zmeny v množstve dôkazov, požadovanej miere istoty a sile výberu ovplyvňovali najmä rýchlosť zániku neúspešných metód, nie to, ktoré metódy napokon prežili. Štandardná aktualizácia pravdepodobnosti a dve metódy založené na vysvetľovacej sile neprežili do poslednej generácie v žiadnej zo skúmaných podmienok, hoci jedna metóda založená na koherencii často zostávala konkurencieschopná dlhšie. Zistenia preto neurčujú jednu všeobecne najlepšiu metódu. Ukazujú, že úspešnosť závisí od štruktúry dôkazov a že vysvetlenie a koherencia môžu v niektorých prostrediach predstavovať užitočné doplnenie štandardnej revízie presvedčení založenej na pravdepodobnosti. Práca zároveň ukazuje, ako možno filozofické myšlienky o racionálnom usudzovaní previesť do podoby testovateľných počítačových modelov.

Kľúčové slová: revízia presvedčení, neistota, koherencia, evolučné výpočty, pravdepodobnostné usudzovanie

Contents

1	Coherence	3
1.1	Historical Development of Coherence	3
1.1.1	Coherence as a Theory of Truth	3
1.1.2	From Truth to Justification	4
1.1.3	Coherence in Philosophy of Science	4
1.2	Coherence, Explanation and Explanatory Bonus	5
1.3	Speed and Accuracy as Epistemic Demands	8
1.4	Coherence and Explanatory Epistemic Agents in Evolutionary Computation	10
2	Measures	13
2.1	Bayesian Measure	13
2.1.1	Origin and difference between “theorem” and “formula”	13
2.1.2	How the formula follows from definitions	14
2.2	Explanatory Measures in Bayesian Confirmation	14
2.2.1	Empirical Explanatory Bonus in the Bayesian Measure	16
2.2.2	Popper’s Measure of Explanatory Power	17
2.2.3	Good’s Measure of Explanatory Power	17
2.3	Coherence Measures	17
2.3.1	Shogenji’s Measure of Coherence	17
2.3.2	Olsson–Glass Measure of Coherence	19
2.3.3	Hartmann–Trpin Measure of Coherence	21
3	Methodology	25
3.1	Computational Application of the Updating Rules	25
3.1.1	Common Bayesian Core	26
3.1.2	Standard Bayesian Updating	26
3.1.3	Empirical Explanatory Bonus Updating	27
3.1.4	Popper-Based Updating	27
3.1.5	Good-Based Updating	28
3.1.6	Shogenji-Coherence Updating	29

3.1.7	Olsson–Glass-Coherence Updating	30
3.1.8	Hartmann–Trpin-Coherence Updating	31
3.1.9	Distinction Between Weighting and Bonus Rules	31
3.1.10	Normalization and Posterior Trajectories	33
3.1.11	Summary of the Applied Method	33
3.2	Evolutionary Selection of Belief-Update Rules	34
3.2.1	Methodological rationale	34
3.2.2	Population structure and initialization	34
3.2.3	Selection-only evolutionary model	35
3.2.4	Generational structure	36
3.2.5	Fitness function	36
3.2.6	Stochastic evidence generation	39
3.2.7	Selection and reproduction	39
3.2.8	Parameter design	40
3.2.9	Diagnostic outputs and final outcome measures	43
4	Results and Discussion	45
4.1	Structure of the Analysis	45
4.2	Main Survival Pattern Across All Simulations	46
4.3	Extinction Timing Across Coin Biases	49
4.3.1	Hartmann–Trpin Under Unfavorable Conditions	50
4.4	Parameter Sensitivity Analysis	51
4.4.1	Effect of Sequence Length	53
4.4.2	Effect of Threshold	54
4.4.3	Effect of Elitism	55
4.5	Special Case: Fully Biased Coin	56
4.6	Summary of Empirical Findings	58
5	General Discussion	59
5.1	Interpretation of the Main Pattern	59
5.2	Relation to Explanation and Coherence	60
5.3	Methodological Limitations	61
5.4	Future Work	62
5.5	Concluding Interpretation	63
	Appendix A: Contents of the Electronic Appendix	75

List of Figures

3.1	Structure of the empirical explanatory bonus updating rule.	27
3.2	Structure of the Popper-based updating rule.	28
3.3	Structure of the Good-based updating rule.	29
3.4	Structure of the Shogenji-coherence updating rule.	30
3.5	Structure of the Olsson–Glass-coherence updating rule.	30
3.6	Structure of the Hartmann–Trpin-coherence updating rule.	31
4.1	Mean final population composition by true coin bias	48
4.2	Bonus-dominant survival trajectory	48
4.3	Coherence-dominant survival trajectory	49
4.4	Mean extinction generation under unfavorable survival regimes	51
4.5	Schematic reconstructions of average population development under sequence-length variation	53
4.6	Threshold sensitivity by survival regime	55
4.7	Elitism sensitivity by survival regime	56
4.8	Comparison of the fully biased coin condition with other Bonus-dominant biases	57

List of Tables

3.1	Computational distinction between continuous weighting and selection-based bonus rules.	32
3.2	Fixed parameters used in the main simulations.	40
3.3	Varied parameters and baseline values used in the simulation experiments.	41
3.4	Experimental configurations used in the parameter-sensitivity analysis.	43
4.1	Mean final agent count by true coin bias across all parameter settings. .	46
4.2	Overall extinction timing by agent group across all simulations.	49
4.3	Extinction timing under unfavorable survival regimes.	50
4.4	Latest average extinction among losing groups under parameter variation.	52
4.5	Extinction timing in the fully biased coin condition compared with other Bonus-dominant biases.	57

Introduction

Human reasoning often takes place under uncertainty. In science, medicine, crisis management, and everyday decision-making, people must revise their beliefs while evidence is incomplete, noisy, and still developing. In such situations, the important question is not only whether a method eventually reaches the correct conclusion, but also whether it reaches it quickly and accurately enough to be useful. This thesis addresses this problem by comparing different belief-update rules in a controlled computational simulation.

The thesis focuses on Bayesian updating, explanatory reasoning, coherence-based reasoning, and evolutionary selection. Its theoretical starting point is the concept of coherence, which is important in epistemology and philosophy of science. Coherence describes the way beliefs, hypotheses, or pieces of evidence support one another as parts of a larger system. A coherent set is not only free of contradiction; its parts also fit together and form an intelligible whole. This is especially relevant in science, where theories are usually evaluated as networks of hypotheses, observations, assumptions, and explanatory relations rather than as isolated claims.

Modern formal epistemology has attempted to make coherence more precise through mathematical measures. This thesis works with Shogenji's measure, the Olsson–Glass measure, and the Hartmann–Trpin measure. These approaches allow coherence to be treated not only as a philosophical idea, but also as a computational criterion that can be implemented and tested. The thesis also examines explanation in belief revision. Standard Bayesian updating provides a rigorous model of probability revision, but inference to the best explanation suggests that a hypothesis may be valuable not only because it predicts evidence, but also because it explains why the evidence occurs and how it fits into a broader pattern.

The main objective of the thesis is to investigate how different belief-update rules perform when treated as competing epistemic strategies. The aim is not to prove that one rule is universally superior, but to examine which rules survive, dominate, or disappear under repeated uncertain evidence. The simulation is based on repeated coin-toss evidence, where agents update their beliefs about the true bias of a coin.

Their performance is evaluated through a fitness function that combines speed and accuracy: agents are rewarded for reaching high confidence in the true hypothesis

quickly, but they are also penalized for inaccurate probability distributions.

The interdisciplinary character of the thesis is central to its significance. Cognitive science is fundamentally interdisciplinary, and this thesis reflects that by combining several perspectives. From psychology, it draws on the study of heuristics and decision-making under uncertainty. From computer science, it uses computational modeling and evolutionary computation. From mathematics, it relies on probability theory, Bayesian updating, coherence measures, explanatory measures, normalization, and scoring rules. From philosophy, especially epistemology and philosophy of science, it takes questions concerning justification, explanation, coherence, and rational belief revision. Finally, it touches biology through the use of an evolutionary selection framework, where rule groups are evaluated through survival, extinction, and differential success across generations.

The thesis is structured as follows. Chapter 1 introduces coherence, its relation to explanation, and the importance of speed and accuracy in reasoning, while also identifying the research gap that this study aims to fill. Chapter 2 presents the main mathematical tools, including Bayes's theorem, explanatory measures, and coherence measures. Chapter 3 describes the simulation methodology, including the update rules, evidence environment, evolutionary selection model, fitness function, and experimental parameters.

Overall, this thesis shows how philosophical theories of coherence and explanation can be transformed into computationally testable models of belief revision. Its contribution lies in placing Bayesian, explanatory, and coherence-sensitive update rules into a shared evolutionary framework and evaluating their ability to track the true hypothesis quickly and accurately.

Chapter 1

Coherence

1.1 Historical Development of Coherence

Coherence is commonly understood as the idea that beliefs, claims, or pieces of evidence gain strength when they “hang together” as part of an organized whole. A coherent set is not merely a set without contradictions. Its parts must also support, explain, or fit with one another in a meaningful way. This idea has played an important role in epistemology and philosophy of science because both fields often deal with systems of beliefs rather than isolated statements. Historically, coherence developed from a theory of truth, into a theory of justification, and later into a concept that philosophers attempted to compare and measure more formally.

1.1.1 Coherence as a Theory of Truth

The coherence theory of truth is usually associated with late nineteenth- and early twentieth-century idealist philosophy. Bradley (1893) and Joachim (1906) are among the most important figures in this early stage. They rejected the idea that truth could be understood by looking only at separate propositions. Instead, they argued that truth depends on a proposition’s place within a larger, unified system. For Joachim (1906), a judgment is true when it belongs to a complete and internally connected whole. In this sense, truth is not only about agreement with a single fact, but about systematic fit.

This early view is important because it introduced the central image of coherence: beliefs or propositions are stronger when they belong to an integrated structure. However, this strong version of coherence faced serious criticism. Russell (1907) argued that internal coherence alone cannot be sufficient for truth, because different and incompatible systems might each be internally coherent. If two systems both “hang together” but contradict each other, coherence by itself cannot show which one is true.

Because of this objection, many later philosophers became cautious about treating

coherence as a full theory of truth. The basic insight was preserved, but the focus changed. Coherence came to be seen less as a definition of truth and more as a condition for rational belief or justification.

1.1.2 From Truth to Justification

In twentieth-century epistemology, coherence became especially important as a theory of justification. The main question was no longer simply “What is truth?” but “When is a belief justified?” Coherentists argued that a belief is justified when it fits well within a person’s wider system of beliefs. This view developed partly as a response to foundationalism. Foundationalists claim that justification must rest on basic beliefs. Coherentists deny that justification needs this kind of fixed foundation. Instead, they argue that beliefs can support one another within a system.

Blanshard (1939) and Rescher (1973) helped develop this more epistemological understanding of coherence. They emphasized that coherence requires more than logical consistency. A set of beliefs may contain no contradiction and still be weak if its members are unrelated. For Rescher (1973), a genuinely coherent system has organization, connection, and mutual support.

Laurence BonJour’s *The Structure of Empirical Knowledge* is one of the most influential modern defenses of epistemic coherentism. BonJour (1985) argued that empirical beliefs are justified through their role in a coherent system that accounts for experience and contains strong inferential connections. Lehrer (1990) also developed a coherentist approach by focusing on how beliefs fit within a person’s accepted system of reasons, objections, and responses.

Despite these developments, coherentism continued to face objections. Critics argued that a belief system might be internally coherent but still disconnected from reality. Others argued that mutual support among beliefs may be circular. These problems pushed coherentists to explain more carefully what kind of “fit” matters and whether coherence must be connected to experience, reliability, or truth.

1.1.3 Coherence in Philosophy of Science

Coherence is also important in philosophy of science because scientific theories are not usually accepted or rejected one statement at a time. A theory includes hypotheses, observations, models, background assumptions, and methods. Scientists often prefer theories that connect these elements in a systematic way. A theory is more persuasive when it explains many facts together, reduces anomalies, and fits with established knowledge.

W. V. O. Quine’s and J. S. Ullian’s idea of knowledge as a *The Web of Belief* was influential in this context. Quine and Ullian (1970) argued that beliefs are tested

as parts of an interconnected system rather than as completely isolated statements. This helped philosophers understand scientific knowledge as holistic. When evidence challenges a theory, scientists may revise different parts of the system rather than abandon one statement immediately.

This scientific use of coherence also led naturally toward formal approaches. If one theory can be more coherent than another, then philosophers need to explain what “more coherent” means. Lewis (1946) raised an important early issue by asking how agreement among independent sources can increase credibility. Later, Shogenji (1999) made this question central in formal epistemology by asking whether coherence is truth-conducive. His work helped shift the discussion toward probabilistic accounts of coherence.

After Shogenji, philosophers developed and criticized different measures of coherence. The details of these measures belong to later sections. For the present overview, the important point is that coherence became a graded concept. Philosophers no longer asked only whether a set of beliefs is coherent. They also asked whether one set is more coherent than another and under what conditions coherence supports truth or justification.

This formal turn does not replace the older philosophical idea. Rather, it continues it in a more precise way. The original idea was that information becomes stronger when its parts hang together. Contemporary work tries to explain how that relation can be represented, compared, and sometimes quantified. Therefore, the history of coherence can be seen as a movement from systematic unity, to epistemic justification, to the formal analysis of information structure.

1.2 Coherence, Explanation and Explanatory Bonus

The concepts of coherence and explanation are central to the methodological logic of this study. The simulation does not treat belief updating as a process in which hypotheses are evaluated only by the likelihood of each new observation. Several of the update rules introduce an additional preference for hypotheses that fit the evidence especially well. In the model, this preference is implemented through a bonus mechanism. The purpose of this section is therefore to justify that mechanism philosophically. The bonus is interpreted as an operationalized form of explanatory preference, while the coherence measures determine, in some rules, which hypothesis qualifies for that preference.

The connection between explanation and belief revision is most clearly expressed in the tradition of inference to the best explanation. Harman (1965) argued that nondeductive reasoning is often better understood as inference to the hypothesis that

best explains the evidence, rather than as simple enumerative induction. On this view, evidence supports a hypothesis not merely because the hypothesis is compatible with the evidence, but because the hypothesis explains why that evidence occurs. Harman did not formulate this idea as a computational bonus rule, nor did he present a formal theory of coherence. Nevertheless, his account provides the philosophical basis for the central methodological idea used here: a hypothesis may receive additional epistemic support when it is judged to provide the best explanation of the evidence. This point is important because the non-standard update rules in the simulation are not meant to reject Bayesian updating entirely. The Bayesian update remains the common foundation. However, some rules add an extra explanatory component after the Bayesian step. This reflects the idea that likelihood alone may not capture all the epistemic force of an explanation. In this sense, the bonus mechanism represents a simplified computational version of inference to the best explanation: after ordinary probabilistic updating, an additional advantage is given to the hypothesis that satisfies a specified explanatory or coherence-based criterion.

BonJour's coherentist epistemology helps explain why coherence can function as such a criterion. BonJour (1985) argues that empirical justification depends on the role of beliefs within a coherent system. Coherence is not merely the absence of contradiction. It includes relations of mutual support, integration, inferential connection, and the ability to resolve anomalies. Explanation is especially important within such a system because explanatory relations can unify otherwise separate beliefs and show why different pieces of evidence belong together. At the same time, BonJour does not reduce coherence entirely to explanation. Coherence is broader than explanatory fit alone. This distinction is useful for the present methodology: the model does not assume that coherence and explanation are identical, but it treats coherence as one important way in which explanatory adequacy can be identified.

Thagard's theory of explanatory coherence makes the relationship between explanation, coherence, and computation more explicit. Thagard (1989) models explanatory coherence as a process in which hypotheses and evidence are evaluated through networks of positive and negative constraints. Hypotheses cohere with evidence when they explain it, when they are consistent with other accepted propositions, and when they contribute to an integrated explanatory structure. Contradictions and unresolved conflicts reduce coherence. Thagard's model is not the same as the update rules implemented in this study. The present simulation does not reproduce ECHO or its connectionist architecture. However, Thagard's work is methodologically important because it shows that explanatory fit and coherence can be treated as computationally evaluable properties. This supports the idea that a model can use coherence to guide the selection of a preferred explanatory hypothesis. The present study builds on this general idea in a deliberately simplified way. In the coherence-based update rules,

the coherence measure does not directly determine the size of the bonus. Instead, it determines which hypothesis receives the bonus. This is a crucial methodological distinction. The coherence score functions as a selection criterion: it identifies the hypothesis that makes the observed evidence sequence hang together most strongly according to a particular formal measure. The bonus then expresses an explanatory preference for that selected hypothesis. Thus, the coherence-based rules merge coherence and explanation in a specific operational sense: coherence is used to identify the best explanatory candidate, and the bonus represents the additional support given to that candidate.

Douven's work is especially relevant to this interpretation because it does not only discuss the relation between Bayesian updating and inference to the best explanation, but also gives this relation a concrete computational form. In his EXPL-style rule, the hypothesis judged to provide the best explanation receives an additional bonus, and his coin-toss examples provide an important methodological predecessor for the simulations reported here. Standard Bayesianism treats rational updating as governed by conditionalization, whereas IBE allows explanatory considerations to influence belief revision. Douven (2013) argues that major Bayesian objections to IBE, including Dutch Book and inaccuracy-minimization arguments, do not decisively show that IBE is irrational. His work supports the methodological possibility that explanatory considerations can legitimately supplement Bayesian updating. For the present study, this is important because the simulation compares pure Bayesian updating with rules in which explanatory or coherence-based considerations modify the posterior distribution.

Douven and Wenmackers (2017) extend this issue into social epistemology by comparing Bayesian and IBE-based reasoners in a simulation model. Their results suggest that, under the assumptions of their model, IBE-style updating can perform well in group settings. The relevance of this work is not that it directly validates the exact bonus rules used in this thesis. Rather, it shows that explanatory updating can be formally modeled and compared with Bayesian updating in a controlled computational environment. This supports the broader methodological strategy of the present study: different update rules can be treated as competing epistemic strategies and evaluated through simulation.

Harris and Hahn (2009) provide a further reason not to treat coherence and Bayesian reasoning as simple opposites. Their study of multiple testimonies shows that coherence among reports affects how believable those reports are, and that this effect can be modeled within a Bayesian framework. This is relevant because the present simulation also keeps Bayesian updating as the shared core. Coherence-sensitive rules are not presented as abandoning probability theory. Rather, they use coherence to decide whether one hypothesis should receive additional explanatory support after the Bayesian update has already been applied.

Taken together, this literature motivates the central methodological move of the thesis. Explanation provides the general rationale for giving extra support to a hypothesis that does more than merely fit the latest observation. Coherence provides one way of identifying which hypothesis has this stronger explanatory status, especially when the evidence is considered as an accumulated pattern rather than as isolated data points. The simulation operationalizes this relationship through the bonus mechanism. In the empirical bonus rule, the bonus is assigned to the hypothesis closest to the observed frequency. In the coherence-based rules, the bonus is assigned to the hypothesis that scores highest on the relevant coherence measure. In both cases, the bonus represents an IBE-style preference for the hypothesis that best organizes or accounts for the evidence. This also explains why the simulation distinguishes between continuous explanatory weighting and selection-based bonus updating. The Popper-based and Good-based rules use explanatory measures as continuous modifiers: every hypothesis is adjusted according to its explanatory score. By contrast, the new update rules based on the Shogenji, Olsson–Glass, and Hartmann–Trpin coherence measures use those measures as selection devices: the formula identifies the most coherent hypothesis or hypotheses, and the fixed bonus is then assigned to them. This design avoids treating different coherence measures as if their numerical values were directly comparable across scales. Instead, coherence determines the target of explanatory preference, while the externally fixed bonus controls the strength of that preference. In this explicit sense, the simulation merges explanation and coherence: explanation supplies the epistemic role of the bonus, while coherence supplies the criterion by which the bonus is assigned.

1.3 Speed and Accuracy as Epistemic Demands

Scientific reasoning is often evaluated in terms of accuracy, reliability, and evidential justification. However, in many real-world contexts, accuracy alone is not sufficient. In emergencies, pandemics, medical triage, disaster response, and other time-pressured situations, a decision may lose much of its practical value if it is reached too late. The central problem is therefore not only whether an inference eventually becomes correct, but whether it becomes correct quickly enough to guide effective action. In this sense, timeliness is not external to good scientific reasoning. It is part of what makes reasoning useful under urgent conditions.

The COVID-19 pandemic illustrates this point clearly. Hsiang et al. (2020) estimate that large-scale anti-contagion policies substantially reduced the growth of COVID-19 infections across several countries. Pei et al. (2020) similarly show that intervention timing had major consequences in the United States: implementing control measures

even one week earlier would have substantially reduced infections and deaths. These findings show that, under exponential spread, delay is not a neutral waiting period. It changes the scale of the problem itself. A decision that is more accurate but arrives too late may therefore be less valuable than a faster decision that is sufficiently reliable and open to later revision. The same logic appears in urgent medical decision-making. In acute myocardial infarction, De Luca et al. (2004) found that treatment delay in primary angioplasty was associated with increased one-year mortality, with each 30-minute delay increasing relative risk. In sepsis care, Seymour et al. (2017) found that delays in completing emergency treatment bundles and administering antibiotics were associated with higher mortality. These cases show that speed is not merely a practical convenience. In time-sensitive contexts, speed becomes part of the effectiveness of the intervention itself. This does not mean that science should become careless or that fast decisions should replace accurate decisions. Rather, it shows that scientific and medical reasoning often involve a speed–accuracy trade-off. Sequential decision-making research models decisions as processes in which evidence accumulates over time until a decision threshold or boundary is reached (Ratcliff & McKoon, 2008; Wald, 1945). Higher thresholds generally require more evidence and may improve caution, but they also delay action. Lower thresholds allow faster decisions, but they may increase the risk of error. Drugowitsch et al. (2019) and Liesefeld and Janczyk (2019) emphasize that speed and accuracy must therefore be evaluated together rather than separately. A method that is fast only because it is reckless is not desirable, but a method that is accurate only after excessive delay may also be inadequate in urgent contexts.

This creates a strong motivation for novel approaches that combine speed and accuracy in scientific inference. Belief-updating rules should not be judged only by their final posterior state after all evidence has been observed, but also by the trajectory through which they reach that state. A rule may eventually identify the true hypothesis but still be weak if it requires too much evidence before becoming useful; conversely, a rule may converge quickly but remain unreliable if it overreacts to early evidence and produces poorly calibrated probabilities. This is the motivation behind the present study, which compares several belief-update rules in a controlled probabilistic environment according to how quickly and accurately they move toward the true hypothesis. The fitness function directly operationalizes this aim through two components: time-to-threshold, which rewards early high confidence in the true hypothesis, and a Brier-score penalty, which evaluates probabilistic accuracy across the evidence sequence (Brier, 1950; Gneiting & Raftery, 2007). Together, these components prevent the model from rewarding speed alone or accuracy alone. The evolutionary-computation framework is therefore suitable because it creates a selection environment in which fixed update-rule groups survive only if they achieve a favorable balance between rapid convergence and calibrated accuracy. In this way, the study evaluates belief revision dynamically, as

a process of inquiry rather than only as a final result, and connects explanation- and coherence-sensitive updating to the broader demand for fast but disciplined scientific reasoning.

1.4 Coherence and Explanatory Epistemic Agents in Evolutionary Computation

Recent work in formal epistemology has increasingly moved from static questions about whether belief systems are coherent or truth-conducive toward computational models in which epistemic rules are tested dynamically. This shift is especially relevant for evolutionary-computation approaches, where update rules can be treated as competing epistemic strategies and evaluated according to performance over repeated evidence. For the present project, the most important literature is not the social-communication branch of agent-based epistemology, but rather the work that operationalizes Bayesian, explanatory, and coherence-sensitive rules as computational procedures. In this respect, the project follows a growing methodological trend: epistemic principles are no longer assessed only by abstract normative argument, but also by their behavior in simulated environments.

Douven’s work provides one of the clearest precedents for evaluating epistemic rules computationally. Douven (2019) uses agent-based optimization, described as a form of evolutionary computing, to compare combinations of epistemic principles according to how well they balance speed and accuracy. Accuracy is measured by average Brier penalty, while speed is measured by how quickly agents reach high confidence in the true hypothesis. Although Douven’s model includes social interaction, the methodological structure is highly relevant here because it treats epistemic rules as candidates whose performance can be compared under repeated evidence and selection-like pressure. The present project adopts this general computational inspiration while removing the social-communication component: agents do not exchange beliefs, form peer groups, or aggregate others’ credences. Instead, rule types compete only through their individual success in updating on stochastic evidence. Douven’s later work also supports this ecological view of updating. Douven (2025) argues that non-Bayesian updating rules may outperform Bayesian updating in some environments and proposes reinforcement learning as a way to choose update rules when environmental conditions are unknown. This strengthens the motivation for comparing fixed update rules experimentally rather than assuming that Bayesian conditionalization is always optimal.

The explanatory-agent component of the present project is most directly connected to Douven and Schupbach’s explanationist alternative to Bayesianism. Douven and Schupbach (2015) argue that probabilistic updating need not be purely Bayesian and

that explanatory considerations can improve models of belief revision. Their work is especially useful because it discusses formal measures of explanatory power associated with Popper and Good. Popper-style and Good-style measures compare how expected the evidence is under a hypothesis with how expected it is given the prior predictive distribution. This makes them suitable for implementation as continuous explanatory weights, because each hypothesis receives a graded explanatory score rather than simply being selected as the single best explanation. This supports the methodological distinction in the present simulation between explanatory rules, which modify the whole posterior distribution continuously, and coherence rules, which are used as selection devices for assigning a fixed bonus.

The coherence side of the literature begins with the formal debate about whether coherence is truth-conducive. Shogenji (1999) introduced an influential probabilistic treatment of coherence, while Olsson (2002) sharpened the coherence-and-truth problem by asking under what conditions coherent information should be expected to be more reliable. Glass (2007) brought this discussion closer to inference to the best explanation by proposing that coherence measures can be used to rank competing explanations according to how well they cohere with evidence. This is directly relevant to the present project because the coherence rules are not treated as independent alternatives to evidential updating; rather, they guide which hypothesis receives additional support after the Bayesian update. The literature also shows that coherence is not a single unified quantity. Different measures, such as Shogenji’s measure, the Olsson–Glass relative-overlap measure, and later Hartmann–Trpin-style measures, behave differently and are not necessarily on the same numerical scale. For this reason, the present project’s use of coherence as a ranking or selection criterion, rather than as a direct multiplicative weight, is methodologically defensible.

The current state of the art is represented especially by recent simulation studies that make coherence operational (Baccini et al., 2026; Trpin & Justin, 2025). Baccini et al. (2026) explicitly move beyond static truth-conduciveness questions and introduce a dynamic iterative setting in which an agent repeatedly receives noisy signals and uses a coherence measure to decide which signals to trust or discard. Their simulations compare several coherence measures in this active role. A central finding is that, when signals are not too noisy, agents using the Olsson–Glass relative-overlap measure outperform agents using the other tested measures; however, all measures perform worse as the signal environment becomes increasingly degraded. This is perhaps the closest current analogue to the present project because coherence is treated not merely as a passive property of an information set, but as an active heuristic guiding repeated learning.

Another highly relevant contribution is Trpin and Justin (2025) that simulates agents in Bayesian-network environments who use coherence to filter evidence under

varying levels of noise and bias. In some versions of their model, drops in coherence are treated as higher-order evidence that something has gone epistemically wrong. Their results show that coherence-based filtering can improve accuracy in noisy environments, but that it can also mislead agents when evidence is systematically biased. This conditional result is important for the present project because it suggests that coherence-sensitive updating should not be assumed to be universally superior. Its success depends on the structure of the evidential environment. The present evolutionary design is well suited to testing this point because coherence agents survive only if their rule performs well relative to Bayesian and explanatory competitors across different stochastic evidence conditions.

Taken together, this literature shows that the present project occupies a relatively novel position. Douven (2019, 2025) demonstrate that epistemic rules can be evaluated by speed, accuracy, and evolutionary-style performance, but his central models usually include social learning or adaptive rule choice. Recent coherence studies show that coherence can function as an active heuristic for filtering or selecting information, but they do not usually place coherence rules in direct evolutionary competition with Bayesian and explanatory update rules (Baccini et al., 2026; Trpin & Justin, 2025). The originality of the present research therefore lies in combining these two directions: it treats Bayesian, explanatory, and coherence-sensitive update rules as fixed epistemic agent types within a selection-only evolutionary algorithm. Agents do not communicate, mutate, or recombine; instead, their survival depends on how efficiently and accurately their update rule tracks the true hypothesis under stochastic coin-toss evidence. This makes the project novel in three main respects: it constructs new update rules from formal coherence measures, places those coherence-sensitive rules in direct evolutionary competition with Bayesian and explanatory update rules, and evaluates coherence not only by belief accuracy but also by the survival and extinction patterns of rule groups.

Chapter 2

Measures

2.1 Bayesian Measure

Bayes's theorem describes how to update probabilistic beliefs about hypotheses when new evidence arrives. It lies at the conceptual heart of Bayesian inference and is written in its familiar school form as

$$P(H | E) = \frac{P(E | H) \cdot P(H)}{P(E)} \quad (2.1)$$

where H is a hypothesis (a “cause”), E an observed event (the “evidence”), $P(H)$ the prior probability of H , $P(E | H)$ the likelihood of observing E given H is true, and $P(H | E)$ the posterior probability of H after observing E . When there is a partition of possible causes $H_1, \dots, H_n$ for event E the normalized (multi-hypothesis) form is

$$P(H_i | E) = \frac{P(E | H_i) \cdot P(H_i)}{P(E)} \quad (2.2)$$

which can be obtained by the law of total probability, i.e.

$$P(E) = \sum_{i=1}^n P(E \cap H_i) \quad (2.3)$$

2.1.1 Origin and difference between “theorem” and “formula”

Historically, the idea behind Bayes's theorem predates the compact formula above. Bayes (1763) published a posthumous essay that solved a specific inverse probability problem (a binomial inference). Bayes's work established the basic idea of inferring causes from observed events, but he did not express the result in the general symbolic conditional-probability notation used today. Laplace (1774/1986) independently generalized and widely applied the idea in the late 18th and early 19th centuries; he stated the inverse-probability principle in words and applied it to many problems (urns, celestial mechanics, and parameter estimation). However, the neat algebraic identity above

— written with explicit conditional-probability notation — only becomes natural and commonplace after probability is given an axiomatic measure-theoretic foundation and conditional probability is formally defined. Kolmogorov (1950) sets out the axioms and the definition of conditional probability from which Bayes’s identity follows immediately as an algebraic consequence. Thus, historically: Bayes supplied the motivating special case, Laplace gave the general verbal rule and large applications, and Kolmogorov provided the modern axiomatic framework that makes the compact formula a straightforward consequence of definitions.

2.1.2 How the formula follows from definitions

Under standard measure-probability notation, conditional probability is defined (when $P(E) > 0$) by

$$P(H \mid E) = \frac{P(H \cap E)}{P(E)} \quad (2.4)$$

Equivalently, $P(H \cap E) = P(H)P(E \mid H) = P(E)P(H \mid E)$

Solving for $P(H \mid E)$ yields the schoolbook Bayes formula: $P(H \mid E) = P(E \mid H)P(H)/P(E)$. The normalization denominator $P(E)$ can be expanded by conditioning on a partition of hypotheses to give the multi-hypothesis form above.

According to Jaynes (2003), rational belief revision follows a precise rule: when evidence E is observed, one’s updated (posterior) belief in a hypothesis H should be proportional to how well H predicts E , weighted by the prior belief in H , i.e. $P_{\text{new}}(H) = P_{\text{old}}(H \mid E)$. This formula ensures coherence, meaning beliefs remain internally consistent and aligned with the rules of probability, avoiding contradictions or arbitrary changes. In practice, this framework underpins fields ranging from statistics and machine learning to scientific inference, where hypotheses are continuously tested and refined as data accumulates. By encoding both prior knowledge and new observations in a mathematically rigorous way, probability theory provides a normative standard for rational learning from experience.

2.2 Explanatory Measures in Bayesian Confirmation

Explanatory value (or explanatory merit) is the idea that some hypotheses are better explanations of observed data than others. In everyday terms: two stories might both predict the same fact, but one of them actually makes the fact make sense — it shows how the fact follows from underlying mechanisms, regularities, or laws. In philosophy of science this feature is treated as a scientific virtue: we prefer hypotheses that unify diverse facts, reduce surprise, or show how a phenomenon arises from prior structure. Explanatory value is usually taken to be a graded property of hypothesis–evidence

pairs (how much the evidence counts in favor of the hypothesis as an explanation). Philosophers and formal epistemologists have tried to (a) characterize what makes a hypothesis explanatory, and (b) measure that explanatory quality in probabilistic terms so it can play a role in rational belief change.

Peirce (1931–1958) introduced **abduction** as the reasoning move that generates explanatory hypotheses: abduction asks, “what, if true, would best explain this surprising fact?”. Concretely, abduction is the creative step that proposes candidate explanations; it is distinct from deduction (deriving consequences that should follow if a hypothesis is true) and from induction or confirmation (the subsequent process of testing and revising credences in light of data). In scientific practice, abduction supplies hypotheses worth testing, while probabilistic updating (e.g., Bayesian conditioning) is the normative mechanism for assessing and comparing those candidates.

Hempel and Oppenheim (1948) formalize one influential account of explanation: the covering-law, or deductive-nomological, model. On this view, a phenomenon is explained when it can be deduced from general laws together with particular conditions. Hempel (1945) connected explanation and confirmation by noting that evidence confirms a hypothesis when the hypothesis makes that evidence expected or entailed. When logic is taken literally, however, this commitment yields counterintuitive consequences - the best known being the ravens paradox. The paradox arises because “All ravens are black” is logically equivalent to its contrapositive “All non-black things are non-ravens.” If observing a black raven counts as confirming the universal, then, by logical equivalence, observing a non-black non-raven (say, a green apple) will also be treated as confirming the universal. Intuitively the apple feels irrelevant. Bayesian and other probabilistic treatments show that the degree of confirmation will typically be much smaller for the apple than for a black raven (because of relative reference classes), but the paradox highlights that an entailment-based account of confirmation can miss the notion of relevance that explanatory judgments rely on.

Bromberger (1966) raises the asymmetry of explanation problem, which targets the idea that explanation should run in only one direction: if A explains B , then B should not count as explaining A , even when the same lawlike connection could be used to derive either statement. Hempel’s covering-law account seems to miss this point because it treats explanation as a matter of deducing the explanandum from general laws plus initial conditions, which implies that, in some cases, the derivation can be reversed and still appear equally legitimate. Bromberger’s point is that this creates a problem for why-questions: a satisfactory explanation must answer the question actually asked, not merely provide a valid deduction. Thus, although the facts may be symmetrically connected by the same laws, our explanatory practices are not symmetric; for example, knowing the height of a flagpole and the angle of the sun can explain the length of its shadow, but not vice versa, even though both can be derived from the same

geometry. The deeper lesson is that explanation is not merely logical derivation but a relation shaped by causal, directional, and question-relative features that Hempel’s model cannot fully capture.

These worries motivated attempts to capture explanatory merit quantitatively. Popper (1959) and Good (1960) proposed early probabilistic measures that tie explanation to likelihood raising and information: Popper sketched a normalized difference measure of how much more expected evidence E is under hypothesis H than in general, and Good developed the idea of weight of evidence in an information-theoretic, log-odds form. Later, van Fraassen (1989) gave a prominent Bayesian-context critique of simply adding extra weight to the best explanation, arguing that naive explanatory-bonus updates can violate diachronic coherence and conflict with standard conditionalization. Douven (1999) subsequently formalized specific bonus update rules, often called IBE* variants, and analyzed their properties and costs. More recent work by Schupbach and Sprenger (2011) and Hartmann and Trpin (2026b) develops probabilistic and coherence-based measures of explanatory power that attempt to capture explanatory virtues while preserving probabilistic coherence.

2.2.1 Empirical Explanatory Bonus in the Bayesian Measure

Under ordinary Bayesian conditionalization, the posterior probability is proportional to the prior multiplied by the likelihood: $P(H \mid E) \propto P(H) \cdot P(E \mid H)$. The idea behind an explanatory bonus is to modify this rule so that the hypothesis that best explains the evidence receives extra weight before normalization. Informally, some formalizations of inference to the best explanation (IBE)—a rule favoring hypotheses with greater explanatory power—proceed by adding a non-negative bonus term $b(H, E)$ to the Bayesian numerator and subsequently renormalizing across competing hypotheses. Douven (2016) and Douven and Wenmackers (2017) discuss such explanationist alternatives to standard Bayesian updating:

$$P_{\text{IBE}}(H_i \mid E) = \frac{P(E \mid H_i)P(H_i) + b(H_i, E)}{\sum_j (P(E \mid H_j)P(H_j) + b(H_j, E))}.$$

When all $b(H, E) = 0$, this reduces to ordinary conditionalization; when a positive bonus is assigned to one hypothesis, it receives extra posterior weight relative to pure Bayes. van Fraassen (1989) emphasized that even conceptually natural variants of this move can produce diachronic incoherence, violating some of the standard constraints that motivate Bayesian updating. (Douven, 1999) and later authors therefore studied precise formulations, including how to choose b , whether it should be additive or multiplicative, and how large it may be without problematic consequences. Later work by Schupbach and Sprenger (2011) and Hartmann and Trpin (2026b) develops probabilistic measures of explanatory power that aim to respect Bayesian norms.

2.2.2 Popper's Measure of Explanatory Power

In *The Logic of Scientific Discovery*, Popper (1959) introduced a specific measure of explanatory power. He proposed that a hypothesis H explains evidence E to the extent that E would be more expected if H is true. Popper's measure of explanatory power can be written as

$$\text{EP}_{\text{Popper}}(E, H) = \frac{P(E | H) - P(E)}{P(E | H) + P(E)}.$$

This ratio ranges from -1 to $+1$ and increases as $P(E | H)$ grows larger relative to the prior probability $P(E)$. In other words, it is the difference between how likely E is under H and how likely E is in general, normalized by their sum. As Schupbach and Sprenger (2011) note, this is Popper's measure of explanatory power. They also observe that Popper's measure is ordinally equivalent to Good's measure and to a posterior odds ratio. Thus, Popper's work provides an early explicit probabilistic formula for explanatory power.

2.2.3 Good's Measure of Explanatory Power

Building on Popper's idea, Good (1960) developed the notion of weight of evidence as a measure of how much data support a hypothesis. Good's treatment in "Weight of Evidence, Corroboration, Explanatory Power, Information and the Utility of Experiments" framed explanation in information-theoretic terms. In one common form, Good's measure can be written as

$$\text{EP}_{\text{Good}}(E, H) = \log \frac{P(H | E)}{P(H)} = \log \frac{P(E | H)}{P(E)}.$$

This captures how much more supported H becomes after observing E , or equivalently how much more expected E is under H than in general. As Schupbach and Sprenger (2011) note, Good's measure is closely related to Popper's measure and can be understood as an information-theoretic measure of explanatory power. Today, Good's measure is commonly associated with Bayesian confirmation and log-ratio measures of evidential support.

2.3 Coherence Measures

2.3.1 Shogenji's Measure of Coherence

Shogenji (1999) proposed an influential quantitative definition of epistemic coherence. His central idea was that coherence among a set of beliefs can be understood in probabilistic terms as the degree to which those beliefs hang together, that is, the extent to which the truth of one belief increases the likelihood of the truth of the others.

Shogenji's insight was to treat coherence as a matter of probabilistic dependence: beliefs are coherent when they are more likely to be true together than they would be if they were independent.

Thus, $C(B_1, B_2) > 1$ indicates positive dependence, meaning that the beliefs cohere; $C(B_1, B_2) = 1$ indicates probabilistic independence, or neutral coherence; and $C(B_1, B_2) < 1$ indicates negative dependence, meaning that the beliefs are incoherent. This ratio therefore measures the degree of deviation from probabilistic independence, essentially giving a correlational index of how strongly the beliefs support or undermine each other.

Following Shogenji (1999), the pairwise Shogenji measure of coherence is expressed as

$$C(B_1, B_2) = \frac{P(B_1 \cap B_2)}{P(B_1)P(B_2)}.$$

Here, the numerator $P(B_1 \cap B_2)$ represents the joint probability that both beliefs are true together, while the denominator $P(B_1)P(B_2)$ represents the probability they would have if they were independent. Thus, $C(B_1, B_2) > 1$ indicates positive dependence, $C(B_1, B_2) = 1$ indicates independence, and $C(B_1, B_2) < 1$ indicates negative dependence. This ratio therefore measures the degree of deviation from probabilistic independence, essentially giving a correlational index of how strongly the beliefs support each other.

Shogenji generalizes this approach to sets of three or more beliefs $\{B_1, B_2, \dots, B_N\}$. The general formula is:

$$C(B_1, \dots, B_N) = \frac{P(B_1 \cap \dots \cap B_N)}{P(B_1) \dots P(B_N)}.$$

This measure quantifies how much more likely the entire set of beliefs is to be true together than it would be if each belief were true independently of the others. The calculation is straightforward once the relevant probabilities are known: compute the joint probability of all beliefs being true, compute the product of their individual probabilities, and then divide the former by the latter. The resulting ratio provides a single numerical degree of coherence for the set.

Conceptually, Shogenji's coherence measure captures the idea of probabilistic interdependence: coherence is a function of how much the truth of the beliefs raises or lowers the joint probability of the set. His measure treats independence as the baseline of neutrality and evaluates coherence as the degree of positive correlation among beliefs. This definition connects epistemic coherence to the probabilistic structure of evidence and belief, offering a rigorous tool for comparing how well sets of propositions fit together.

2.3.2 Olsson–Glass Measure of Coherence

The Olsson–Glass measure of coherence is a probabilistic index designed to quantify the extent to which two propositions overlap in the probability mass assigned to them. Formally, for propositions A and B , the measure is defined as:

$$C(A, B) = \frac{P(A \wedge B)}{P(A \vee B)}$$

This ratio, often called the relative-overlap measure, ranges from 0 to 1. A value of 0 indicates complete incompatibility, meaning that there is no joint support and $P(A \wedge B) = 0$. A value of 1 indicates perfect alignment, meaning that all probability assigned to either A or B lies in their intersection, so that $P(A \vee B) = P(A \wedge B)$. The numerator captures joint support, while the denominator normalizes that support by the total probability mass committed to at least one of the propositions, thereby producing an index of coherence.

Olsson (2002) and Glass (2002) independently introduced this measure in 2002. In both cases, they proposed essentially the same formulation in different contexts: Olsson in a philosophical analysis of coherence and truth, and Glass in the study of Bayesian networks and explanation. Its intuitive appeal lies in its direct probabilistic interpretation: coherence is the fraction of the relevant probability mass that is mutually supporting, which makes the index easy to compute when joint and marginal probabilities are available. Importantly,

$$P(A \vee B) = P(A) + P(B) - P(A \wedge B)$$

that is, the sum of the marginal probabilities minus the joint probability, since the overlap between A and B would otherwise be counted twice.

Several formal properties of the relative-overlap measure are noteworthy. First, the measure is symmetric:

$$C(A, B) = C(B, A)$$

by the commutativity of $\wedge$ and $\vee$. This is desirable because coherence should not depend on the order in which propositions are considered. For example, the coherence of “the sky is cloudy” and “it will rain today” should be the same regardless of order. Second, monotonicity with respect to strengthening of propositions does not hold unconditionally: making A stronger, for example by replacing A with $A \wedge X$, can either increase or decrease C , depending on how the added conjunct interacts probabilistically with B . Third, the measure respects probabilistic independence in the expected way. For independent events with small marginal probabilities, the numerator $P(A \wedge B)$ is small relative to $P(A \vee B)$, yielding low coherence. Conversely, highly correlated or mutually entailing propositions produce high coherence values.

The relative-overlap index has several practical advantages for applied work. Its bounded $[0, 1]$ scale facilitates interpretation and comparison across different pairs of propositions, and the normalization by $P(A \vee B)$ avoids spurious inflation of coherence in cases where both events have small but jointly overlapping probability. Moreover, because it is expressed directly in terms of joint and marginal probabilities, it adapts naturally to Bayesian frameworks and probabilistic graphical models, which allow straightforward computation.

However, the measure is not without limitations. It is sensitive to the choice of the probability distribution: different background information or priors can change $P(A \wedge B)$ and $P(A \vee B)$ and thus the coherence score. Moreover, the measure evaluates only pairwise overlap and does not capture important structural aspects of coherence in larger networks of beliefs, where considerations such as explanatory power, evidential support, or logical entailment may be relevant. Extensions and alternatives have therefore been proposed that attempt to capture multi-propositional coherence or incorporate normative considerations.

These limitations have been discussed in subsequent work on coherence theory, where it has been argued that simple overlap measures may fail to capture explanatory or inferential relations between propositions. In particular, broader critiques of probabilistic coherence approaches emphasize that coherence alone may not guarantee truth-conduciveness or epistemic justification. Olsson (2005) develops this criticism in detail.

A further limitation concerns the behavior of the relative-overlap measure with respect to expanding sets of propositions. Increasing the size of a set cannot increase coherence under this measure. For example, consider the set containing the propositions “this animal cannot fly” and “this animal is a bird,” which appears relatively incoherent. Extending this to a second set containing “this animal cannot fly,” “this animal is a bird,” and “this animal is a penguin” intuitively increases coherence, since the additional proposition explains the apparent tension. However, probabilistically,

$$P(A \wedge B \wedge C) \leq P(A \wedge B)$$

and

$$P(A \vee B \vee C) \geq P(A \vee B)$$

so the coherence value cannot increase. This type of counterexample is closely related to well-known cases in the literature, such as the Tweety or penguin examples, which highlight the inability of simple relative-overlap measures to capture explanatory coherence. Koscholke et al. (2019) discusses this type of limitation in relation to relative-overlap measures of coherence. In short, “non-flying animal + bird” appears incoherent, whereas “non-flying animal + bird + penguin” appears more coherent because the penguin proposition explains the non-flying bird.

In sum, the Olsson–Glass relative-overlap measure occupies a useful niche: it is a compact, interpretable, and computationally tractable index of pairwise probabilistic coherence, well suited to formal analyses within Bayesian epistemology and probabilistic modeling, but its scope is intentionally limited to quantifying overlap rather than prescribing broader epistemic virtues.

2.3.3 Hartmann–Trpin Measure of Coherence

In their paper “Coherence of Information: What It Is and Why It Matters,” Hartmann and Trpin (2023) introduced the Hartmann–Trpin coherence measure. In that paper, they argue that earlier probabilistic measures each captured only part of the picture. The Shogenji measure, introduced by Shogenji (1999), is built around probabilistic dependence, so it does well at tracking correlation and deviation from independence. The Olsson–Glass measure, introduced by Olsson (2002) and Glass (2002), is by contrast an overlap measure: it asks how likely it is that all items in the set are jointly true relative to at least one being true. Hartmann and Trpin’s motivation was to combine these two ideas into a single measure that keeps the strength of both while avoiding the limitations of each taken alone.

The original paper presents the Hartmann–Trpin measure, usually written as $\text{Coh}_{\text{OG}+}$, as a ratio of two overlap scores: the actual overlap of the information set divided by the overlap that would obtain if the same propositions were probabilistically independent but kept the same marginal probabilities. In the original notation, the formula is

$$\text{Coh}_{\text{OG}+}(S) := \frac{P(H_1, \dots, H_n)}{P(H_1 \vee \dots \vee H_n)} \bigg/ \frac{\tilde{P}(H_1, \dots, H_n)}{\tilde{P}(H_1 \vee \dots \vee H_n)}$$

Read in words, this says: first compute how tightly the propositions hang together in the actual distribution; then compute the same thing in a counterfactual independence baseline; then compare the two. The tilde distribution, $\tilde{P}$, is the benchmark distribution with the same individual probabilities as the real one but with all propositions treated as independent.

This is also why the measure is meant to satisfy both of the classical desiderata in coherence theory: Dependence and Agreement. Dependence captures the idea that coherent propositions should be correlated rather than probabilistically separate. Agreement captures the idea that they should jointly fit together, or overlap, rather than merely being statistically linked in some abstract way. Shogenji’s measure is strong on Dependence because it directly tracks departure from independence, but it does not build in the overlap intuition as explicitly. Olsson–Glass is strong on overlap and Agreement, but by itself it cannot recover the dependence-based notion of absolute coherence. Hartmann and Trpin explicitly present $\text{Coh}_{\text{OG}+}$ as a compromise that combines these two motivations.

This construction is what lets the measure do more than Shogenji or Olsson–Glass alone. If the propositions are positively correlated, the actual relative overlap is larger than the independence-based relative overlap, so the measure comes out above 1. If they are independent, the two overlaps match and the measure is 1. If they are negatively correlated, it falls below 1. Hartmann and Trpin (2023) report this clean dependence behavior for information pairs and triples, while also noting that Dependence does not hold in general for larger sets. That limitation is not treated as a bug in their presentation; rather, it is part of the tradeoff they are willing to accept in order to keep the measure useful for small sets while still preserving the overlap intuition.

The 2024 paper, “A New Posterior Probability-Based Measure of Coherence,” keeps the same core idea but shifts the probabilistic frame. The key novelty is that most earlier Bayesian coherence measures are prior-based, whereas this proposal is posterior-based. So the core idea does not change, but the perspective does. Instead of asking only how propositions cohere from the prior point of view, the newer paper asks how coherence looks after updating on evidence. That makes the measure more directly tied to belief revision and posterior assessment, while preserving the original aim of comparing actual coherence with what would be expected under independence (Hartmann & Trpin, 2024).

Trpin and Justin (2025) takes the idea further by putting coherence into a learning environment. The authors use Bayesian networks to simulate agents under varying levels of noise and bias, and some agents treat drops in coherence as higher-order evidence that something has gone epistemically wrong. The main result is a mixed one: this strategy can improve belief accuracy in noisy environments, but it tends to mislead when evidence is systematically biased. So coherence can work as a useful heuristic for filtering evidence, but only when the noise is not structured in a way that pushes the agent toward the wrong hypothesis.

Baccini et al. (2026) extend this discussion to truth-tracking and data selection. The paper explicitly treats Hartmann–Trpin as one of the main coherence measures in the simulation and notes again that, while the measure behaves well for small sets, it does not satisfy Dependence for larger sets. The study’s broader point is that coherence-guided data selection can support truth-tracking, but the performance profile depends on the measure and on the environment, so no single measure dominates in every regime. That fits the Hartmann–Trpin story nicely: the measure is designed as a principled compromise, not a universal solution. In this sense, $\text{Coh}_{\text{OG}+}$ is not a measure that fully satisfies both classical desiderata. Rather, as Hartmann and Trpin (2026a) emphasize, it should be understood as a measure that is guided by the intuitions behind Dependence and Agreement, even though it does not technically capture either of them completely.

Overall, Hartmann and Trpin’s measure is best understood as a carefully designed

compromise. It was introduced in to bring together the dependence intuition of Shogenji and the overlap intuition of Olsson–Glass. It then got refined in later work, including a posterior-based variant in 2024 and dynamic truth-tracking applications in 2025 and 2026. Its strongest virtue is that it gives a clean, interpretable benchmark: coherence is measured relative to independence, not merely by raw overlap. Its main weakness is also clear: once information sets become large, the measure no longer preserves Dependence in full generality.

Chapter 3

Methodology

3.1 Computational Application of the Updating Rules

This part of the methodology describes how seven update rules were applied in the simulation: standard Bayesian updating, the empirical explanatory bonus rule, henceforth “Bonus” in the figures and tables, Popper-based updating, Good-based updating, and three coherence-measure-based rules using the Shogenji, Olsson–Glass, and Hartmann–Trpin measures. The purpose of this section is not to reintroduce the philosophical or mathematical justification of the formulas, which has already been discussed above, but to explain how those formulas were operationalized in a common computational framework.

The simulation was built around a simple binary evidence environment: repeated coin tosses. Each observation was coded as either heads or tails, where heads was represented by 0 and tails by 1. The model compared competing hypotheses about the bias of the coin. The hypothesis space consisted of eleven possible coin biases:

$$0.0, 0.1, 0.2, \dots, 1.0.$$

Each hypothesis therefore represented one possible value of $P(\textit{heads})$. For example, the hypothesis $H = 0.7$ represents a coin that produces heads with probability 0.7 and tails with probability 0.3. The likelihood structure was constructed directly from these hypotheses: if the observation was heads, the likelihood of that observation under a given hypothesis was equal to the hypothesis value itself; if the observation was tails, the likelihood was one minus that value.

All update rules were applied to the same initial prior distribution, the same likelihood structure, and the same observed sequence. This was done so that differences between the posterior distributions could be attributed to the update rule itself rather than to differences in the input data. In the baseline setting, the initial prior assigned equal probability to all eleven hypotheses, meaning that the model began without favoring any particular coin bias.

The updating process was sequential. After each observed toss, the current prior was transformed into a posterior distribution. That posterior then became the prior for the next observation. As a result, the model did not merely calculate a final belief state after the whole sequence; it recorded the entire trajectory of belief revision. For a sequence of n observations, each rule therefore produced $n+1$ probability distributions: the initial prior plus one posterior after each toss.

3.1.1 Common Bayesian Core

All seven rules share the same basic inferential structure: they revise beliefs by comparing how well each hypothesis predicts the observed evidence. In the standard Bayesian rule, the prior probability of each hypothesis is multiplied by the likelihood of the current observation under that hypothesis, and the Bayesian denominator normalizes the posterior distribution. In the implementation, this is equivalent to dividing the resulting values by their sum. For the modified rules, however, normalization is needed again after the explanatory weighting or additive bonus has been applied, because these extra adjustments can change the total probability mass.

This Bayesian rule serves as the baseline. It contains no explanatory or coherence-based modification. It simply asks: given the current observation, which hypotheses made this observation more likely?

The remaining six rules modify this Bayesian baseline in different ways. However, the nature of the modification is not the same in every case. This distinction is important methodologically. Some rules apply a fixed additive bonus to selected hypotheses, whereas others use the explanatory formula to continuously adjust all hypotheses. Therefore, it would be misleading to describe all non-standard rules as simply “adding a bonus.”

The update rules can be divided into three groups:

1. **Pure Bayesian updating**, where no extra adjustment is made.
2. **Continuous explanatory weighting**, where the explanatory measure modifies the Bayesian update for every hypothesis.
3. **Selection-based additive bonus updating**, where a formula identifies the best hypothesis or hypotheses, and a fixed bonus is then added to them.

3.1.2 Standard Bayesian Updating

This updating rule is based on the Bayesian measure discussed in Section 2.1. The standard Bayesian rule was used as the control condition. At each step, the rule updated the prior by multiplying each hypothesis by the likelihood of the observed toss.

If heads was observed, hypotheses with higher heads probabilities received more posterior weight. If tails was observed, hypotheses with lower heads probabilities received more posterior weight.

No additional explanatory or coherence-based term was introduced. This rule therefore represents ordinary conditionalization and provides the reference point for evaluating the behavior of the other six rules.

3.1.3 Empirical Explanatory Bonus Updating

This updating rule is based on empirical explanatory bonus introduced in Section 2.2.1. The empirical explanatory bonus rule introduces a fixed additive bonus after the ordinary Bayesian update. The rule first calculates the running proportion of heads in the sequence observed so far. For example, after the sequence heads, heads, tails, heads, the running heads ratio is $3/4 = 0.75$. The rule then identifies which hypothesis, or hypotheses, are closest to this observed heads ratio. If the running ratio is 0.75, the closest hypotheses in the eleven-point grid are typically 0.7 and 0.8. The fixed bonus is then divided equally among the closest hypotheses. After the bonus is added, the distribution is normalized again. This rule therefore uses the observed frequency of heads as a direct guide for selecting which hypotheses receive extra support. The size of the bonus is not determined by the distance between the hypothesis and the observed ratio. Rather, the distance only determines which hypothesis or hypotheses receive the fixed bonus.

The structure of the empirical explanatory bonus updating is shown in Figure 3.1.

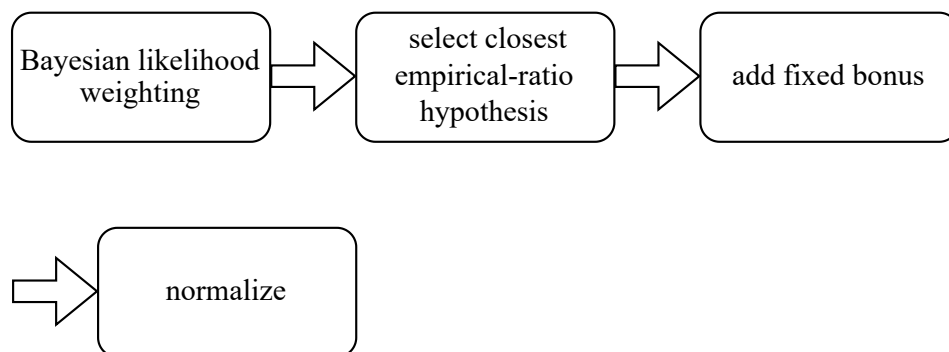


Figure 3.1: Structure of the empirical explanatory bonus updating rule.

3.1.4 Popper-Based Updating

This updating rule is based on Popper's measure of explanatory power, discussed in Section 2.2.2. The Popper-based rule does not add a fixed bonus to a selected hypoth-

esis. Instead, Popper’s explanatory measure is used as a continuous adjustment to the Bayesian update for every hypothesis.

At each observation, the model calculates how likely the observed event is under each hypothesis and compares this with the overall expected probability of that event under the current prior. A hypothesis receives a positive Popper adjustment when it makes the evidence more expected than the model expected overall. It receives a negative adjustment when it makes the evidence less expected than the model expected overall. The crucial implementation point is that the Popper score does not merely choose a winning hypothesis. It modifies the update across the whole hypothesis space. Each hypothesis receives an adjustment whose size depends on its Popper explanatory score and on the strength of the Bayesian term itself.

The structure of the Popper-based rule is shown in Figure 3.2.

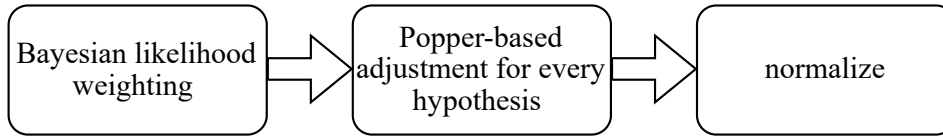


Figure 3.2: Structure of the Popper-based updating rule.

This means that Popper-based updating remains Bayesian in its core, because the update still begins from prior multiplied by likelihood. However, it is not pure Bayesian conditionalization, because the posterior is further shaped by an explanatory weighting.

3.1.5 Good-Based Updating

This updating rule is based on Good’s measure of explanatory power, discussed in Section 2.2.3. The Good-based rule is implemented in the same general class as the Popper-based rule. It also does not assign a fixed bonus to a selected hypothesis. Instead, Good’s explanatory measure is used to continuously scale the Bayesian update for each hypothesis. The rule compares the probability of the observed evidence under each hypothesis with the overall probability of the evidence under the current prior. Because Good’s measure is logarithmic, it can produce stronger positive or negative adjustments than Popper’s bounded measure. Hypotheses that make the evidence substantially more expected than the prior predictive probability receive increased weight; hypotheses that make the evidence less expected receive reduced weight. In practical implementation, this rule requires numerical safeguards because logarithmic calculations are sensitive to zero probabilities. Very small values are therefore used to prevent undefined logarithms or invalid posterior weights. After the Good-based adjustment is applied, the posterior is normalized.

The structure of the Good-based rule is shown in Figure 3.3.

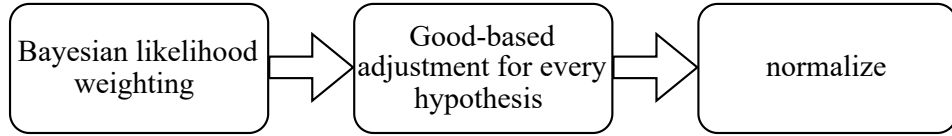


Figure 3.3: Structure of the Good-based updating rule.

As with the Popper rule, the Good rule is Bayesian in its core but not identical to ordinary Bayes. The explanatory measure changes the strength of support assigned to each hypothesis before normalization.

3.1.6 Shogenji-Coherence Updating

This updating rule is based on Shogenji's measure of coherence, introduced in Section 2.3.1. The Shogenji-based rule belongs to the selection-based bonus. This means that Shogenji's coherence measure is not used to scale every hypothesis proportionally. Instead, it is used to identify the hypothesis or hypotheses with the highest coherence score. A fixed bonus is then added to those hypotheses. This use of a coherence measure as a criterion for assigning an explanatory bonus follows the methodological rationale developed in Section 1.2. After each observation, the model first performs the ordinary Bayesian update. It then evaluates the observed sequence up to that point. The coherence calculation asks which hypothesis makes the accumulated evidence hang together most strongly relative to an independence baseline. In the coin-toss setting, this means that the rule does not evaluate only the most recent toss, but the partial sequence observed so far. For example, after the sequence heads, heads, tails, the coherence calculation evaluates which coin-bias hypothesis makes that three-toss pattern most coherent. The hypothesis with the highest Shogenji coherence score receives the fixed bonus. If several hypotheses tie for the highest coherence score, the bonus is divided equally among them. The resulting distribution is then normalized. The important methodological point is that the Shogenji measure determines the recipient of the bonus, not the size of the bonus. A hypothesis with a much higher coherence score does not receive a larger bonus than a hypothesis that only narrowly wins. The magnitude of the bonus is fixed externally.

The structure of the Shogenji-coherence rule is shown in Figure 3.4.

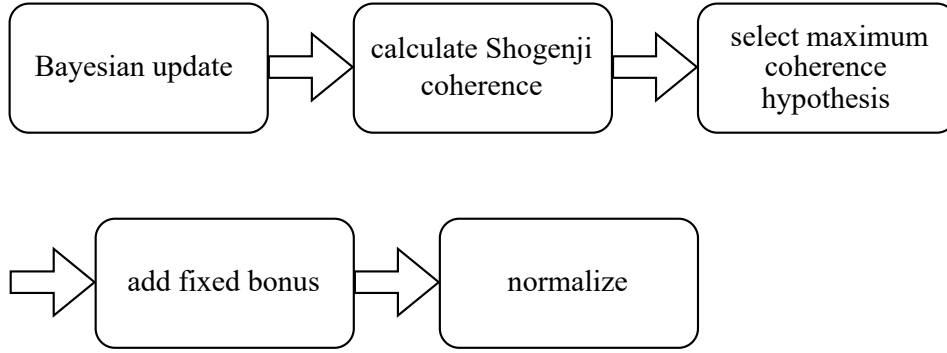


Figure 3.4: Structure of the Shogenji-coherence updating rule.

3.1.7 Olsson–Glass-Coherence Updating

This updating rule is based on Olsson–Glass measure of coherence, introduced in Section 2.3.2. The Olsson–Glass rule is implemented in the same selection-based way as the Shogenji rule. The model first performs an ordinary Bayesian update for the current observation. It then evaluates the observed subsequence using the Olsson–Glass relative-overlap idea. In this application, the observed sequence is treated as the evidence pattern to be explained by each hypothesis. The model compares the support each hypothesis gives to the actual sequence with a union-like measure of relevant support. This allows the rule to identify which hypothesis gives the strongest overlap with the observed evidence pattern. Once the Glass–Olsson coherence scores are calculated, the hypothesis or hypotheses with the highest score are selected. The fixed bonus is then added only to those selected hypotheses and the result is normalized again. As with the Shogenji rule, the Olsson–Glass score is used as a ranking or selection device. It does not determine the amount of additional probability mass assigned. The bonus magnitude remains fixed; the formula only determines where that bonus is applied. This use of a coherence measure as a selection criterion for an explanatory bonus follows the methodological rationale developed in Section 1.2.

The structure of the Glass–Olsson-coherence rule is shown in Figure 3.5.

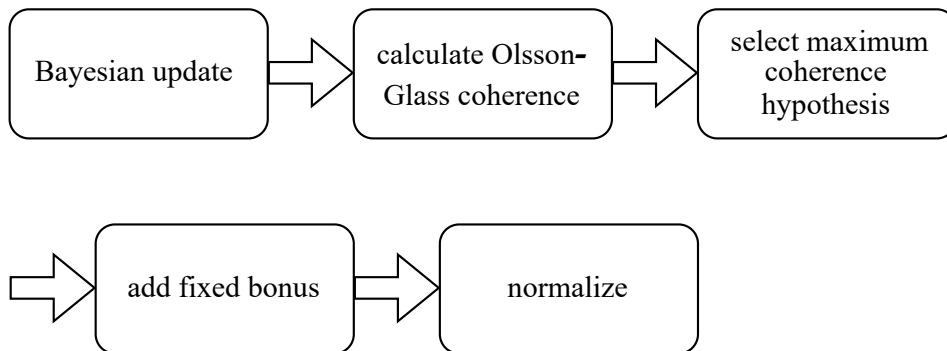


Figure 3.5: Structure of the Olsson–Glass-coherence updating rule.

3.1.8 Hartmann–Trpin-Coherence Updating

This updating rule is based on Hartmann–Trpin measure of coherence, introduced in Section 2.3.3. The Hartmann–Trpin rule is also implemented as a selection-based additive bonus rule. The model first updates the prior by ordinary Bayesian conditionalization. It then evaluates the observed sequence using a Hartmann–Trpin-style coherence comparison. The Hartmann–Trpin calculation compares two quantities. First, it evaluates how strongly each hypothesis actually overlaps with the observed evidence pattern. Second, it evaluates how much overlap would be expected under an independence benchmark, where the same marginal probabilities are preserved but the relevant propositions are treated as independent. The coherence score is then based on the comparison between actual overlap and independence-based overlap. In the implementation, this produces one Hartmann–Trpin coherence score for each hypothesis. The hypothesis or hypotheses with the highest score are selected as the most coherent with the evidence sequence observed so far. The fixed bonus is then added to those selected hypotheses and the distribution is normalized again.

As with the Shogenji and Olsson–Glass rules, the Hartmann–Trpin score determines which hypothesis receives the bonus. It does not determine the size of the bonus. This is an important distinction because the implementation does not reward hypotheses in proportion to their coherence scores. It uses the coherence measure to identify the best-supported candidate and then applies a fixed additive adjustment. The philosophical reason for combining coherence-based selection with a fixed explanatory bonus is discussed in Section 1.2.

The structure of the Hartmann–Trpin-coherence rule is shown in Figure 3.6.

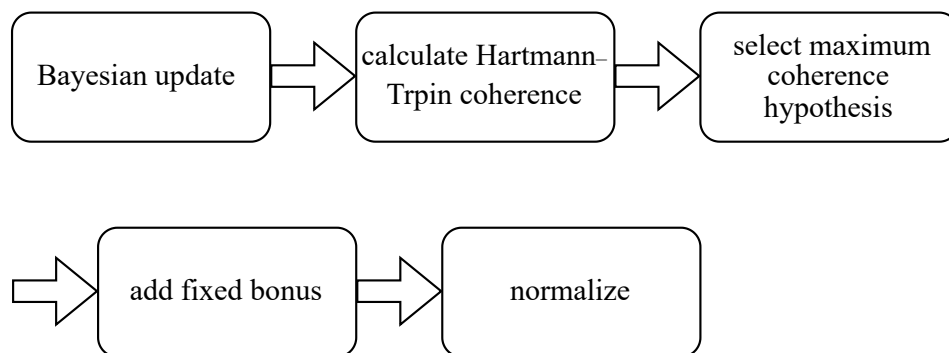


Figure 3.6: Structure of the Hartmann–Trpin-coherence updating rule.

3.1.9 Distinction Between Weighting and Bonus Rules

A central methodological distinction in the implementation is the difference between continuous weighting and selection-based bonus assignment. The Popper and Good

rules use their explanatory measures as continuous modifiers. Every hypothesis is affected by the explanatory score. Hypotheses with positive explanatory scores are strengthened relative to ordinary Bayes, while hypotheses with negative explanatory scores are weakened. The explanatory formula directly affects the magnitude of the adjustment. By contrast, the empirical bonus, Shogenji, Olsson–Glass, and Hartmann–Trpin rules use their criteria to select one or more hypotheses. Once the selected hypotheses are identified, a fixed bonus is added. In these rules, the formula or criterion determines the target of the bonus, not its size. The bonus size is controlled by the externally chosen bonus parameter. The distinction between these two implementation strategies is summarized in Table 3.1.

Table 3.1: Computational distinction between continuous weighting and selection-based bonus rules.

Rule	Type of modification	Role of formula or criterion
Standard Bayes	No modification	Updates only by prior and likelihood.
Empirical bonus	Fixed additive bonus	Selects the hypothesis closest to the running heads ratio.
Popper	Continuous explanatory weighting	Adjusts every hypothesis according to its Popper score.
Good	Continuous explanatory weighting	Adjusts every hypothesis according to its Good score.
Shogenji	Fixed additive bonus	Selects the hypothesis with the highest Shogenji coherence.
Olsson–Glass	Fixed additive bonus	Selects the hypothesis with the highest overlap coherence.
Hartmann–Trpin	Fixed additive bonus	Selects the hypothesis with the highest coherence.

This clarification is important because the seven rules do not all implement explanation or coherence in the same computational way. Some use explanatory measures to reshape the whole posterior distribution, while others use coherence measures to choose which hypothesis receives an externally fixed bonus. Popper’s and Good’s measures naturally return graded explanatory support for every hypothesis relative to the current item of evidence, making them suitable for continuous weighting. The coherence measures, however, evaluate how well the observed evidence sequence hangs together under each hypothesis and are therefore used to identify the most coherent candidate explanation. The fixed bonus then operationalizes an IBE-style preference for the best explanation without allowing differences in the numerical scale of the coherence measures to dominate the update. This is important because Shogenji, Olsson–Glass, and

Hartmann–Trpin scores are not on the same scale and do not have the same numerical interpretation. Using them as direct multiplicative weights would make the comparison sensitive to scale differences rather than to the qualitative role of coherence in selecting a favored hypothesis. Thus, the implemented design separates the role of coherence as a selection criterion from the externally controlled strength of the explanatory bonus.

3.1.10 Normalization and Posterior Trajectories

After every update, the distribution is normalized so that the probabilities over the eleven hypotheses sum to one. This is done after the Bayesian step and, where applicable, after the explanatory or coherence-based modification. Normalization is essential because multiplying by likelihoods, applying explanatory weights, or adding bonuses changes the raw probability values.

Each rule produces a posterior trajectory over time. These trajectories allow comparison of how quickly different rules concentrate probability on particular hypotheses and how strongly they react to different evidence patterns. For example, after repeated heads, all rules tend to move probability mass toward higher-bias hypotheses. However, the degree and speed of this movement differ depending on whether the rule uses pure Bayesian updating, continuous explanatory weighting, or fixed selection-based bonus assignment.

The comparison therefore focuses not only on final posterior probabilities, but also on the dynamic behavior of belief revision across the sequence. This makes it possible to examine whether explanation- or coherence-sensitive rules converge faster, overreact to early evidence, or produce sharper posterior distributions than standard Bayesian updating.

3.1.11 Summary of the Applied Method

In summary, the simulation applies seven update rules to the same binary evidence sequence and the same discrete hypothesis space. The standard Bayesian rule provides the baseline. The empirical bonus rule adds a fixed bonus to hypotheses closest to the observed heads ratio. The Popper and Good rules continuously adjust all hypotheses according to explanatory strength. The Shogenji, Olsson–Glass, and Hartmann–Trpin rules use coherence measures to identify the most coherent hypothesis and then assign a fixed bonus to them. The implementation treats Bayesian updating as the common foundation while varying the way explanation or coherence enters the update process. This allows the analysis to compare pure conditionalization, explanatory weighting, and coherence-guided bonus assignment within a single controlled computational framework.

3.2 Evolutionary Selection of Belief-Update Rules

3.2.1 Methodological rationale

This study uses an evolutionary-computation framework to compare the performance of several fixed belief-update rules under repeated probabilistic inference conditions. Evolutionary computation is generally based on population-level adaptation, where individuals are evaluated according to a fitness function and better-performing individuals are more likely to survive into later generations (Eiben & Smith, 2015; Holland, 1975). In this thesis, however, the evolutionary algorithm is not used to discover new mathematical formulas. Instead, it is used as a controlled selection environment in which predefined epistemic update rules compete against one another. This use of evolutionary computation as a comparative selection environment follows the methodological motivation developed in Section 1.4.

This distinction is important. In many genetic algorithms, individuals may be modified by mutation or recombination. In the present experiment, the individuals represent agents belonging to fixed update-rule groups. The purpose is therefore not to optimize the internal structure of the rules, but to compare how well different rule types survive when exposed to the same stochastic evidence and the same selection pressures. This interpretation is closer to a selection-based or replicator-style model, where population frequencies change because some types perform better than others (Adami et al., 2016; Nowak & Sigmund, 2004). Evolutionary game theory and agent-based evolutionary models commonly use this kind of framework to study how strategies spread, survive, or disappear under selective pressure.

The methodological goal is therefore comparative rather than generative. The algorithm asks: “Given a set of fixed update rules, which rule groups tend to survive, dominate, or die out under evolutionary selection?” This framing makes the evolutionary algorithm appropriate even without mutation and crossover, because the research question concerns the relative success of existing rules, not the invention of new hybrid rules.

3.2.2 Population structure and initialization

The simulation begins with a population of 350 agents. These agents are divided equally among seven update-rule groups, meaning that each rule type starts with 50 agents. This equal initialization is necessary because the study compares the survival of different rule groups. If some groups started with more agents than others, their later survival could partly reflect initial numerical advantage rather than actual performance.

Equal starting conditions are especially important in selection-only models. In such models, population composition changes through differential survival and reproduction.

If a type disappears, it cannot return unless mutation or reintroduction is allowed. Therefore, the initial distribution matters. By assigning exactly 50 agents to each of the seven rule groups, the experiment gives each rule an equal opportunity to reproduce and prevents the results from being biased by unequal initial representation.

The total population size of 350 is also methodologically practical. Evolutionary-computation literature emphasizes that population size is problem-dependent and must balance diversity, reliability, and computational cost (Eiben & Smith, 2015). Here, 350 agents is large enough to give each rule group a meaningful initial presence, while still keeping the simulation computationally feasible across multiple parameter settings and bias values.

3.2.3 Selection-only evolutionary model

A major design choice in this experiment is the exclusion of mutation and crossover. In standard genetic algorithms, mutation introduces random variation, while crossover recombines information from two parents to produce new offspring. These operators are useful when the goal is to search a large solution space and generate new candidate solutions. However, they are not appropriate for the main purpose of this thesis.

The agents in this study are defined by mathematically specified update rules. If mutation or crossover were introduced, the resulting descendants might no longer correspond clearly to any of the original epistemic rules. For example, a mutated or recombined rule could become a hybrid whose philosophical or mathematical interpretation is unclear. This would make the results harder to interpret, because survival would no longer show which original update rule performed best.

For this reason, the algorithm is intentionally designed as a selection-only evolutionary model. Agents reproduce by copying successful rule types, but the rule formulas themselves remain unchanged. This allows population changes to be interpreted directly: if a group survives, it does so because its fixed update rule performs well under the chosen fitness function; if it dies out, it does so because it performs poorly relative to the competing rules.

This design is supported by selection-based evolutionary models and evolutionary game theory, where the central object of study is often the change in frequencies of fixed strategies over time (Adami et al., 2016; Nowak & Sigmund, 2004). In such models, mutation is not always required; indeed, excluding mutation can be methodologically useful when the researcher wants to isolate the effect of selection on predefined types.

Thus, the absence of mutation and crossover is not a weakness of the algorithm. It is a deliberate methodological restriction that preserves the interpretability of the comparison.

3.2.4 Generational structure

The simulation proceeds for 120 generations. Each generation consists of evaluation, selection, and reproduction. First, every agent is evaluated using the fitness function. Second, the best-performing agents are selected. Third, a new population is formed from elite agents and tournament-selected agents.

A fixed generation budget is a standard stopping criterion in evolutionary computation (Eiben & Smith, 2015). In this study, 120 generations are used as a practical observation window. The aim is not to prove theoretical convergence to a global optimum, but to observe whether meaningful survival, dominance, and extinction patterns emerge over time.

The use of generations also makes it possible to track population change dynamically. Rather than only comparing final results, the experiment records how the population composition changes from one generation to the next. These temporal patterns provide more information than a single final population count.

3.2.5 Fitness function

General definition of fitness

The fitness function determines how well an agent performs in the inference task. In this experiment, fitness measures how quickly and accurately an agent assigns high posterior probability to the true hypothesis. Lower fitness values are better.

The fitness function has two components:

$$\text{Fitness} = \text{time-to-threshold} + \text{Brier penalty}.$$

The first term rewards agents that reach a required confidence threshold quickly. The second term penalizes agents whose posterior probability distributions are inaccurate across the sequence. This means that the algorithm rewards both speed and probabilistic accuracy.

This is important because evaluating only speed could reward agents that become confident too quickly even when their posterior distributions are poor. Conversely, evaluating only final accuracy would ignore how efficiently an agent reaches a decision. Combining the two components reflects the speed–accuracy trade-off discussed in sequential decision-making literature. This design directly implements the epistemic motivation developed in Section 1.3, where speed and accuracy are treated as joint demands on good belief revision.

Threshold-based convergence

The first part of the fitness function is based on threshold-based convergence. An agent is considered to have converged when the posterior probability assigned to the true hypothesis reaches or exceeds a specified threshold.

If H_{true} is the true hypothesis and $P_t(H_{\text{true}})$ is the posterior probability assigned to that hypothesis after observing evidence up to time t , convergence occurs when

$$P_t(H_{\text{true}}) \geq \theta.$$

where θ is the confidence threshold.

This threshold-based approach is not arbitrary in principle. In sequential decision theory, decisions are often modeled as processes in which evidence accumulates over time until it reaches a decision boundary (Ratcliff & McKoon, 2008; Wald, 1945). Diffusion decision models, for example, describe decisions as noisy evidence accumulation toward a boundary, where stricter boundaries lead to slower but more cautious decisions (Ratcliff & McKoon, 2008). Bayesian sequential designs also commonly use posterior-probability thresholds as stopping rules (Zhou & Ji, 2024).

In this study, the tested thresholds are:

$$\theta \in \{0.85, 0.90, 0.95\}.$$

These values are best understood as a high-confidence sensitivity band, not as universally correct constants. The threshold of $\theta = 0.85$ represents a moderately strict confidence requirement, $\theta = 0.90$ represents the baseline high-confidence requirement, and $\theta = 0.95$ represents a very strict confidence requirement.

This design allows the experiment to test whether rule survival changes when agents are required to be more or less confident before they are counted as converged. Therefore, the thresholds are not defended as uniquely correct values. Instead, they are defended as a controlled range of confidence levels that operationalize different decision standards.

Convergence speed: time-to-threshold

The main component of fitness is time-to-threshold. This is the number of coin tosses required for the agent to reach the posterior confidence threshold.

For example, if the threshold is $\theta = 0.90$ and an agent assigns probability 0.90 or higher to the true hypothesis after 42 tosses, then its base convergence score is 42. If it reaches the threshold only after 160 tosses, the score is 160. If it never reaches the threshold within the available sequence, it receives a penalty equal to the sequence length plus one.

This means that the algorithm rewards agents that become confidently correct using fewer observations. This is appropriate because the experiment is interested not only in whether a rule eventually identifies the true hypothesis, but also in how efficiently it does so.

The time-to-threshold measure is methodologically aligned with sequential decision-making models, where the timing of a decision is central. A rule that reaches justified confidence earlier has an advantage over a rule that requires much more evidence. At the same time, this speed advantage is moderated by the Brier penalty, which prevents the algorithm from rewarding speed alone.

Brier-score penalty

The second component of the fitness function is the Brier-score penalty. The Brier score is a widely used measure for evaluating probabilistic predictions (Brier, 1950). It is also a proper scoring rule, meaning that it rewards probability assignments that accurately represent uncertainty (Gneiting & Raftery, 2007).

In this experiment, the Brier score compares the agent’s posterior probability vector to an ideal vector. The ideal vector assigns probability 1 to the true hypothesis and probability 0 to all other hypotheses.

If p_j is the posterior probability assigned to hypothesis j , and y_j is the ideal probability for hypothesis j , then the Brier score (BS) is

$$\text{BS} = \frac{1}{K} \sum_{j=1}^K (p_j - y_j)^2$$

where K is the number of hypotheses.

The BS is calculated at each posterior time step and then averaged across the whole evidence sequence. This average Brier value is multiplied by the sequence length so that the penalty is placed on a scale comparable to the time-to-threshold component.

The purpose of this penalty is to prevent misleading fast convergence. Some agents may reach the threshold quickly, but they may do so through unstable or poorly calibrated posterior behavior. For example, an agent might assign high probability to a hypothesis early, but still distribute probability poorly across the rest of the hypothesis space. Without an accuracy penalty, such an agent could be rewarded too strongly simply because it crossed the threshold.

The Brier penalty therefore makes the fitness function stricter. An agent must not only reach high confidence quickly; it must also maintain posterior probabilities that remain close to the true hypothesis across the sequence. This is consistent with the broader methodological point that speed and accuracy should be evaluated together when they can trade off against one another (Drugowitsch et al., 2015; Liesefeld & Janczyk, 2019).

3.2.6 Stochastic evidence generation

Agents are evaluated on generated coin-toss sequences. Each sequence is produced using a specified true coin bias. Since the sequence is stochastic, two runs with the same true bias may still produce different observed evidence patterns.

Repeated stochastic evaluation is necessary because the performance of a rule should not depend on one lucky or unlucky evidence sequence. Randomized algorithms and stochastic simulations should be evaluated through repeated independent runs because chance variation is part of the system being studied (Arcuri & Briand, 2014).

In this experiment, stochastic evaluation serves two purposes. First, it makes the inference task more realistic, because agents must respond to noisy evidence rather than deterministic evidence. Second, it reduces the risk that results are caused by a single unusual sequence.

The experiment also varies the true coin bias. The tested bias values range from 0 to 0.9 with step 0.1. These values allow the experiment to test agent performance across different inferential conditions. A bias of 0.5 represents an unbiased coin, while values farther from 0.5 represent stronger bias. Stronger biases usually provide clearer evidence, while values closer to 0.5 make inference more difficult. This connects to decision-making literature showing that evidence strength affects confidence, decision speed, and accuracy (Drugowitsch et al., 2019; Ratcliff & McKoon, 2008).

3.2.7 Selection and reproduction

Elitism

After all agents are evaluated, a fraction of the best-performing agents is copied directly into the next generation. This mechanism is known as elitism. Elitism is commonly used in evolutionary algorithms to prevent the best individuals from being lost due to random sampling (Eiben & Smith, 2015).

The literature does not support one universally optimal elitism value for all evolutionary algorithms. Instead, elitism is treated as a control parameter whose effect depends on the problem. Some elitism can preserve high-quality individuals, but too much elitism can reduce diversity and increase the risk of premature convergence (Eiben & Smith, 2015; Mills et al., 2015). Mills et al. (2015) also emphasize that evolutionary-algorithm control parameters often require empirical sensitivity analysis rather than fixed universal defaults.

Binary tournament selection

After elites are copied into the next generation, the remaining population slots are filled using binary tournament selection. In binary tournament selection, two agents

are randomly selected from the population. Their fitness values are compared, and the agent with the better fitness value is copied into the next generation. Since lower fitness is better in this experiment, the agent with the lower score wins the tournament. This process is repeated until the new population reaches the required size.

Tournament selection is a standard and well-supported selection method in evolutionary algorithms. It has several advantages. First, it is simple to implement. Second, it does not require fitness values to be normalized. Third, selection pressure can be controlled through tournament size (Blickle & Thiele, 1996; Goldberg & Deb, 1991; Miller & Goldberg, 1995).

Larger tournaments increase selection pressure because stronger individuals are more likely to win. However, strong selection pressure can cause the population to converge too quickly, especially when fitness evaluation is noisy. Miller and Goldberg (1995) specifically discuss tournament selection under noisy fitness conditions, which is relevant here because agents are evaluated on stochastic coin sequences. A binary tournament therefore provides a moderate level of selection pressure. It rewards better-performing agents while avoiding the excessive pressure that could occur with larger tournaments.

3.2.8 Parameter design

Fixed Parameters

Two parameters are fixed across the main simulations and are summarized in Table 3.2.

Table 3.2: Fixed parameters used in the main simulations.

Parameter	Value	Justification
Population size	350	Gives seven rule groups with 50 agents each.
Generations	120	Provides a fixed observation window for survival and extinction dynamics.

The population size of 350 ensures equal representation of all seven rule types at the start of each simulation. The generation count of 120 provides enough time to observe meaningful evolutionary dynamics while keeping the experiment computationally feasible.

Varied Parameters

Three parameters are varied in the simulation experiments: coin sequence length, convergence threshold, and elitism rate. Each parameter is tested at three values, with

Table 3.3: Varied parameters and baseline values used in the simulation experiments.

Parameter	Tested values	Baseline
Coin sequence length	150, 200, 250	200
Convergence threshold	0.85, 0.90, 0.95	0.90
Elitism rate	0.00, 0.05, 0.20	0.05

one value serving as the baseline. The tested values and baseline configuration are summarized in Table 3.3.

The baseline configuration therefore uses a coin sequence length of 200, a convergence threshold of 0.90, and an elitism rate of 0.05.

One-Factor-at-a-Time Parameter Strategy

The experiment uses a one-factor-at-a-time strategy. This means that one parameter is changed while the other two remain fixed at their baseline values.

The full factorial design would require testing every possible combination of sequence length, threshold, and elitism. Since each parameter has three values, a full factorial design would require

$$3 \times 3 \times 3 = 27$$

unique parameter configurations before considering repeated runs and different bias values. This would be computationally expensive.

In agent-based and computational models, OFAT-style sensitivity analysis can be useful as an interpretable starting point for examining how model outcomes change when individual parameters are varied around a baseline configuration (ten Broeke et al., 2016). In this study, the goal is not to estimate a complete response surface, but to conduct a limited local sensitivity analysis. Therefore, this design should not be interpreted as a full factorial design, because OFAT designs do not fully estimate interaction effects between parameters (Czitrom, 1999).

Sequence length configurations For sequence length, threshold and elitism are held at their baseline values. The sequence lengths were chosen based on the amount of evidence available to the agents. A sequence length of 150 represents a shorter evidence condition, where only the fastest-converging agents are expected to reach the threshold. A sequence length of 200 is the baseline because preliminary behavior suggests that most rules can reach convergence around this point. A sequence length of 250 gives slower-converging agents more opportunity to reach the threshold. A large effect of sequence length is not necessarily expected because the fitness function already rewards quick convergence. However, sequence length may still matter for rules that require more evidence before reaching high confidence.

Threshold configurations For threshold, sequence length and elitism are held at baseline. The threshold values represent different levels of required confidence. A threshold of 0.85 allows agents to converge under a less strict confidence requirement. A threshold of 0.90 is the baseline high-confidence condition. A threshold of 0.95 requires very strong posterior confidence. This parameter is especially important because it tests whether rule groups behave differently when the standard for decision-making becomes stricter. Some agents may reach 0.85 quickly but struggle to reach 0.95. Others may converge more slowly but produce more reliable posterior distributions. The Brier penalty helps distinguish between agents that are merely fast and agents that are both fast and probabilistically accurate.

Elitism configurations For elitism, sequence length and threshold are held at baseline. These values allow the experiment to compare no elitism, mild elitism, and strong elitism. This is important because elitism affects how strongly the best agents are protected across generations. With no elitism, strong agents can be lost through stochastic selection. With mild elitism, high-performing agents are protected while diversity is still possible. With strong elitism, successful agents are preserved more aggressively, but weaker groups may disappear faster. The value 0.00 is included as a non-elitist comparison condition. In this case, no agent is automatically preserved, so reproduction depends entirely on tournament selection. The value 0.05 is used as the baseline. This represents mild elitism: the best 5% of agents are preserved, but most of the population is still determined by tournament selection. This gives strong agents some protection while still allowing population change. The value 0.20 represents strong elitism. Here, the top 20% of agents are preserved directly. This condition tests whether stronger elite preservation accelerates dominance or extinction patterns.

Full Experimental Design

Each configuration is evaluated across ten true-bias values. The complete set of parameter configurations is summarized in Table 3.4.

Because each of the seven configurations is evaluated across ten true-bias values, the design produces

$$7 \times 10 = 70$$

total simulations. This design allows the experiment to examine both parameter sensitivity and performance across different evidence-generating conditions.

Table 3.4: Experimental configurations used in the parameter-sensitivity analysis.

Configuration	Sequence length	Threshold	Elitism rate
1	150	0.90	0.05
2	200	0.90	0.05
3	250	0.90	0.05
4	200	0.85	0.05
5	200	0.95	0.05
6	200	0.90	0.00
7	200	0.90	0.20

3.2.9 Diagnostic outputs and final outcome measures

During the simulations, the algorithm produced several intermediate diagnostic outputs. These included average fitness over generations, the best-performing individual within each rule group, and population composition across generations. These outputs were useful for checking whether the evolutionary process behaved as expected during implementation. For example, average fitness helped confirm whether selection was generally favoring agents with better convergence performance. The best-individual output allowed inspection of whether each rule type was capable of producing strong individual performance under some conditions. Population-composition plots helped visually confirm whether the selection process was changing the distribution of agent types over generations.

However, these diagnostic outputs were not used as the main results of the thesis. They were primarily included to validate the functioning of the simulation and to provide a visual check that the evolutionary mechanism was operating correctly. Therefore, the final analysis does not statistically compare average fitness histories, best-individual trajectories, or full population-composition curves.

The main values saved for final analysis were the survival and extinction outcomes of each rule group. A rule group was classified as surviving if at least one agent of that type remained in the population at the end of the simulation. A rule group was classified as extinct if its population count reached zero before the final generation. For extinct groups, the generation at which extinction occurred was recorded.

Thus, the formal outcome measures used in the analysis were:

1. whether each rule group survived until the end of the simulation;
2. how many agents of each surviving group remained at the final generation;
3. if a rule group died out, the generation in which extinction occurred.

This distinction is important because the thesis focuses on evolutionary survival and extinction of rule groups, rather than on detailed fitness-curve analysis. The auxiliary plots and intermediate values served as implementation checks, while the saved survival/extinction data provided the basis for the final comparison between update-rule groups.

Chapter 4

Results and Discussion

4.1 Structure of the Analysis

This chapter presents and discusses the results of the evolutionary simulations. The analysis focuses on survival, extinction, and dominance of the seven agent groups across the tested simulation conditions. The primary outcome is whether each rule group survived until the final generation or became extinct. The secondary outcome is the generation at which extinct groups disappeared from the population. These outcome measures are defined methodologically in Section 3.2.9.

The full simulation design produced 70 simulations: seven parameter configurations evaluated across ten true coin-bias values. The parameter structure behind these simulations is described in Section 3.2.8. Each simulation began with 350 agents, divided equally among seven rule groups, with 50 agents per group. Since the population size remains fixed, the final number of agents in each rule group provides a direct measure of evolutionary dominance.

The analysis is organized in three stages. First, the main survival pattern across all simulations is examined. This is the strongest part of the results because it uses the full dataset and reveals a stable division between Bonus-dominant and coherence-dominant conditions. Second, parameter sensitivity is analyzed for sequence length, confidence threshold, and elitism. These analyses are interpreted more cautiously because the data become smaller once they are divided by parameter value and survival regime. Third, a specific additional pattern is examined: the speed with which Bonus agents reached full dominance in the fully biased coin condition compared with other Bonus-dominant biases.

The results are descriptive rather than inferential. The simulations were not designed as a high-powered statistical experiment with repeated trials under every possible parameter combination. Instead, they form a controlled computational comparison of fixed update-rule groups under selected parameter variations. Therefore, the results

are interpreted as patterns within the tested simulation environment, not as universal statistical laws.

The population-composition plots generated during the simulation are also treated carefully. The actual plots show non-linear population dynamics, including fluctuations in agent-group sizes before final dominance occurs. Because the full set of 70 plots is too large for the main text and because some generated plots contain internal code labels, only selected cleaned visualizations are used in the main chapter. The complete set of original simulation plots is placed in the electronic appendix. Any schematic plots used in the main text are simplified visual reconstructions based on average extinction timing, not as direct records of the actual generation-by-generation trajectories.

4.2 Main Survival Pattern Across All Simulations

The strongest result of the experiment is the emergence of a stable survival pattern across all simulations. Final survival was not randomly distributed among the seven rule groups. Instead, the results separated into two clear regimes. In most coin-bias conditions, Bonus agents became the only surviving group. In a smaller set of coin-bias conditions, the final population was composed entirely of coherence-based agents, namely Shogenji and Olsson–Glass agents.

Table 4.1 reports the mean final number of agents by true coin bias. Values are averaged across the seven parameter configurations and rounded to the nearest whole agent. Since every simulation contains 350 agents, each row sums to 350.

Table 4.1: Mean final agent count by true coin bias across all parameter settings.

Coin bias	Bayes	Bonus	Popper	Good	Shogenji	Olsson–Glass	Hartmann–Trpin
0.0	0	350	0	0	0	0	0
0.1	0	350	0	0	0	0	0
0.2	0	350	0	0	0	0	0
0.3	0	0	0	0	240	110	0
0.4	0	350	0	0	0	0	0
0.5	0	0	0	0	179	171	0
0.6	0	350	0	0	0	0	0
0.7	0	0	0	0	156	194	0
0.8	0	350	0	0	0	0	0
0.9	0	350	0	0	0	0	0

The table shows that Bonus agents dominated the final population for seven out of ten coin-bias values. At coin biases 0.0, 0.1, 0.2, 0.4, 0.6, 0.8, and 0.9, Bonus agents reached the maximum possible final population of 350 agents. In those same conditions, all other rule groups had a final population of zero. This means that Bonus agents were the only surviving group in these conditions.

A different pattern appeared at coin biases 0.3, 0.5, and 0.7. In these three conditions, Bonus agents did not survive to the final generation. Instead, the final population was entirely composed of Shogenji and Olsson–Glass agents. At coin bias 0.3, Shogenji agents had the larger final population, with a mean of 240 agents, while Olsson–Glass agents had a mean of 110 agents. At coin bias 0.5, the two coherence-based groups were nearly evenly balanced, with 179 Shogenji agents and 171 Olsson–Glass agents on average. At coin bias 0.7, Olsson–Glass agents had the larger final population, with 194 agents compared with 156 Shogenji agents.

This result must be stated precisely. It does not mean that both Shogenji and Olsson–Glass survived in every individual simulation at coin biases 0.3, 0.5, and 0.7. The more accurate interpretation is that the final surviving population was always coherence-based in these three conditions. Shogenji agents survived in all coherence-dominant simulations. Olsson–Glass agents survived in most of them and survived in all configurations at coin bias 0.7. In the few cases where Olsson–Glass became extinct, Shogenji occupied the remaining population. Thus, the combined final population of Shogenji and Olsson–Glass agents always summed to 350 in the coherence-dominant regime.

The main survival pattern can therefore be summarized as follows. The Bonus rule dominated under most tested true-bias conditions. The coherence-based rules dominated under the intermediate and unbiased conditions 0.3, 0.5, and 0.7. Bayes, Popper, Good, and Hartmann–Trpin agents did not survive to the final generation under any tested condition. This overall pattern is visualized in Figure 4.1.

The two-regime structure is also visible in the survival trajectories. Figure 4.2 shows a Bonus-dominant condition, where the Bonus group expands and becomes the only surviving rule group. Figure 4.3 shows a coherence-dominant condition, where the final population is divided between Shogenji and Olsson–Glass agents.

Together, Figures 4.1, 4.2, and 4.3 support the conclusion that final survival was strongly organized by true coin bias.

The unbiased condition, 0.5, is especially important because the final population was almost evenly divided between Shogenji and Olsson–Glass agents. This suggests that the coherence-based rules were particularly competitive when the evidence stream was less decisive. The neighboring conditions 0.3 and 0.7 also produced coherence dominance, although the internal balance between the two coherence rules shifted. At 0.3, Shogenji had the larger final population; at 0.7, Olsson–Glass had the larger final population. This asymmetry is surprising, because the 0.3 and 0.7 conditions are mirror cases: a 0.3 heads bias is equivalent to a 0.7 tails bias. The difference should therefore be interpreted cautiously and may reflect stochastic variation in the limited number of runs rather than a stable difference between the two coherence rules.

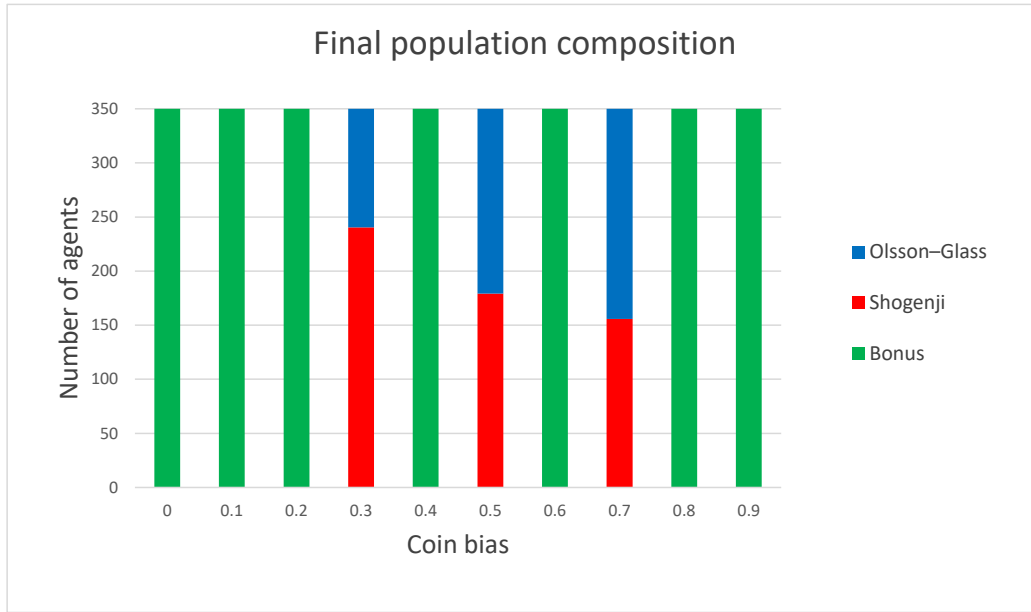


Figure 4.1: Mean final population composition by true coin bias. Bars show the mean final number of agents in each rule group across the seven parameter configurations. Bonus agents dominate at coin biases 0.0, 0.1, 0.2, 0.4, 0.6, 0.8, and 0.9, while Shogenji and Olsson-Glass agents compose the final population at coin biases 0.3, 0.5, and 0.7.

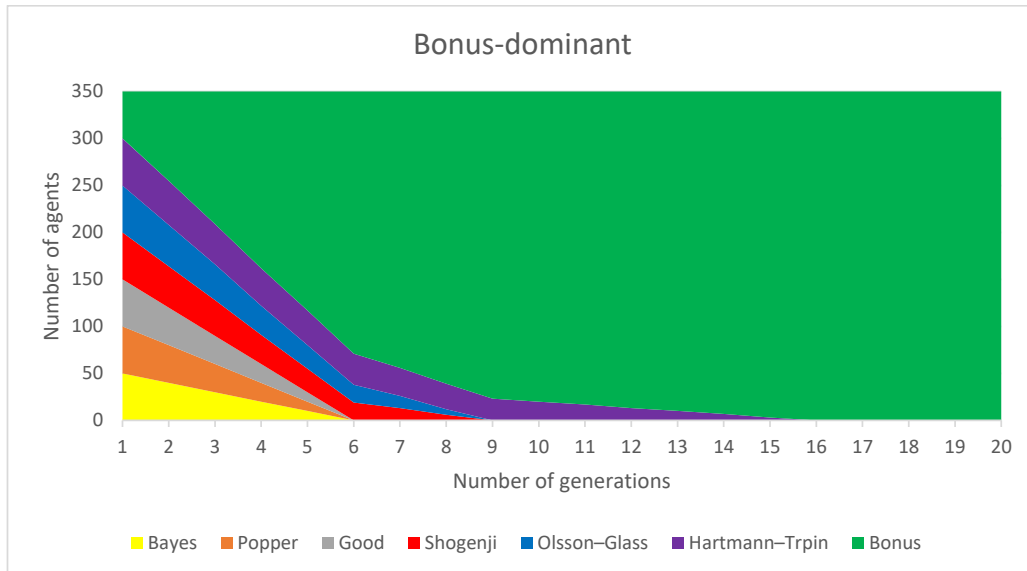


Figure 4.2: Bonus-dominant survival trajectory. The Bonus group expands rapidly and becomes the only surviving rule group, while Bayes, Popper, Good, Shogenji, Olsson-Glass, and Hartmann-Trpin agents disappear before the final generation.

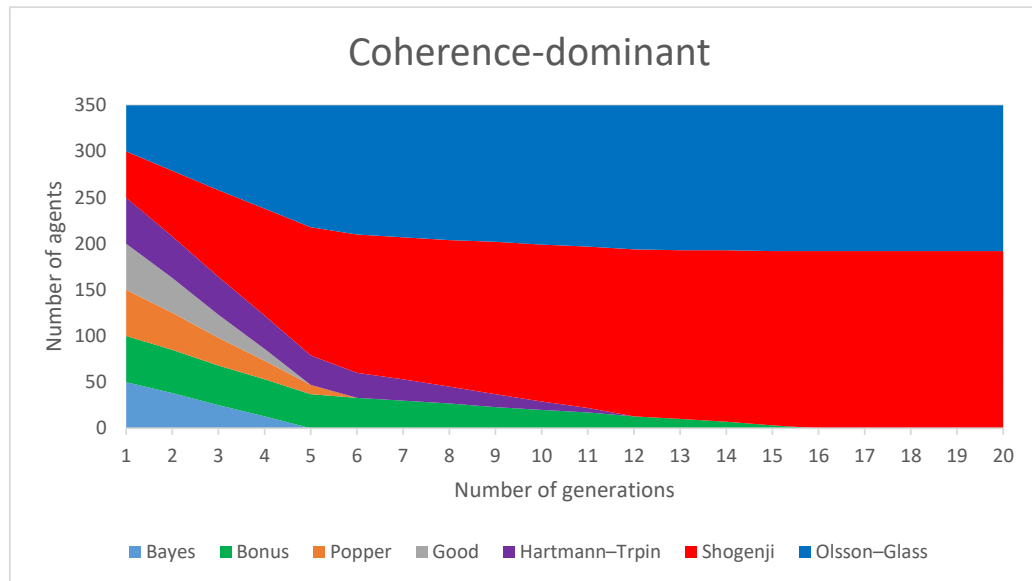


Figure 4.3: Coherence-dominant survival trajectory. The final population is composed of Shogenji and Olsson–Glass agents, while Bayes, Bonus, Popper, Good, and Hartmann–Trpin agents disappear before the final generation.

4.3 Extinction Timing Across Coin Biases

Final population counts show which groups survived, but they do not show how quickly the other groups disappeared. For this reason, extinction generation was also analyzed. Table 4.2 summarizes extinction timing across all simulations.

Table 4.2: Overall extinction timing by agent group across all simulations.

Agent group	Extinction cases	Mean extinction generation	Min	Max
Bayes	70	5	2	10
Bonus	21	16	7	23
Popper	70	6	2	12
Good	70	6	3	11
Shogenji	49	9	3	26
Olsson–Glass	51	12	3	118
Hartmann–Trpin	70	15	4	31

The extinction summary in Table 4.2 shows three important patterns.

First, Bayes, Popper, and Good agents formed an early-extinction cluster. All three groups became extinct in every simulation, and their mean extinction generations were low: 5 for Bayes, 6 for Popper, and 6 for Good. This indicates that these groups were consistently eliminated early by the selection process.

Second, Hartmann–Trpin agents also became extinct in every simulation, but they persisted longer than Bayes, Popper, and Good agents. Hartmann–Trpin agents had a mean extinction generation of 15, with extinction occurring between generations 4 and 31. This means that Hartmann–Trpin agents were never ultimately successful, but they were generally more resistant to early elimination than the other always-extinct groups.

Third, Bonus, Shogenji, and Olsson–Glass agents were regime-dependent. Bonus agents became extinct only in the coherence-dominant regime. Shogenji agents became extinct only in the Bonus-dominant regime. Olsson–Glass agents became extinct in all Bonus-dominant conditions and in two coherence-dominant cases. The maximum extinction generation of 118 for Olsson–Glass reflects a rare case in which it remained competitive until very late in the simulation before disappearing. This pattern matters because extinction timing gives more information than final survival alone. A group that dies at generation 5 and a group that dies at generation 31 are both extinct, but they are not equally weak competitors. Hartmann–Trpin is the clearest example of this distinction. It never survives to the final generation, but it often lasts longer than several other losing groups.

4.3.1 Hartmann–Trpin Under Unfavorable Conditions

Hartmann–Trpin agents occupy an intermediate position in the results. They never became final survivors, but they often persisted longer than some other unsuccessful groups. This becomes clearer when Hartmann–Trpin is compared with other groups under unfavorable conditions.

Table 4.3 compares extinction timing in two regimes. The first is the coherence-dominant regime, where Bonus agents are disadvantaged because Shogenji and Olsson–Glass occupy the final population. The second is the Bonus-dominant regime, where Shogenji and Olsson–Glass are disadvantaged because Bonus agents occupy the final population.

Table 4.3: Extinction timing under unfavorable survival regimes.

Condition	Agent group	Extinction cases	Mean	Min	Max
Coherence-dominant biases	Bonus	21	16	7	23
Coherence-dominant biases	Hartmann–Trpin	21	12	6	18
Bonus-dominant biases	Shogenji	49	9	3	26
Bonus-dominant biases	Olsson–Glass	49	9	3	27
Bonus-dominant biases	Hartmann–Trpin	49	16	4	31

In coherence-dominant conditions, Bonus agents did not survive, but they still

persisted longer than Hartmann–Trpin agents. Bonus agents had a mean extinction generation of 16, compared with 12 for Hartmann–Trpin. This shows that when both groups were disadvantaged by the coherence-dominant environment, Bonus remained competitive for longer.

The comparison changes in Bonus-dominant conditions. In this regime, Shogenji and Olsson–Glass agents became extinct at a mean generation of 9, whereas Hartmann–Trpin agents became extinct at a mean generation of 16. Thus, when Bonus dominated, Hartmann–Trpin agents persisted longer than the coherence agents that eventually died out. This comparison is visualized in Figure 4.4.

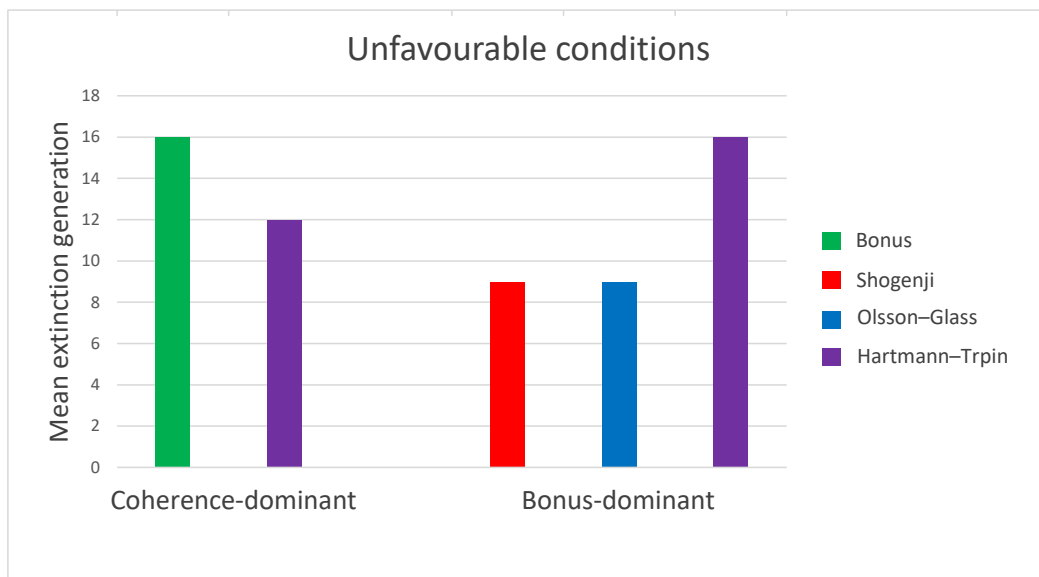


Figure 4.4: Mean extinction generation under unfavorable survival regimes. In coherence-dominant conditions, Bonus agents disappear later than Hartmann–Trpin agents. In Bonus-dominant conditions, Hartmann–Trpin agents persist longer than Shogenji and Olsson–Glass agents, although they still do not survive to the final generation.

This result suggests that Hartmann–Trpin should not be interpreted simply as a weak rule. It never becomes dominant, but it often resists extinction longer than other losing rules. Its performance is therefore best described as persistent but non-dominant.

4.4 Parameter Sensitivity Analysis

The previous section established the main survival pattern across all simulations. This section examines whether the three varied parameters—sequence length, confidence threshold, and elitism—affected extinction timing. The analysis does not treat these parameters as changing final survivor identity, because the identity of the surviving regimes remained stable. Bonus agents continued to dominate in the Bonus-dominant

coin-bias conditions, while Shogenji and Olsson–Glass agents continued to dominate in the coherence-dominant conditions.

The parameter analysis is therefore focused on extinction timing and dominance speed. Specifically, it asks whether losing groups disappeared earlier or later when sequence length, threshold, or elitism changed. These comparisons are descriptive and should be interpreted cautiously. Once the data are divided by parameter value and survival regime, the number of cases in each comparison is small, especially in the coherence-dominant regime. The purpose of this analysis is therefore not to make strong statistical claims, but to assess whether the main survival pattern is robust and whether there are visible timing tendencies.

Table 4.4 summarizes the key parameter-sensitivity results. The reported value is the latest average extinction among losing groups. Lower values indicate that the surviving regime reached dominance earlier on average. In the table, the low column refers to the low, short, or no-elitism parameter setting, while the high column refers to the high, long, or strong parameter setting.

Table 4.4: Latest average extinction among losing groups under parameter variation.

Parameter varied	Regime	Low	Baseline	High	Change	Change
					low	high
Sequence length	Bonus-dominant	17	18	16	-1	-2
Sequence length	Coherence-domin.	17	17	17	0	0
Threshold	Bonus-dominant	17	18	13	-1	-5
Threshold	Coherence-domin.	17	17	18	0	1
Elitism	Bonus-dominant	21	18	13	3	-5
Elitism	Coherence-domin.	17	17	12	0	-5

One additional pattern visible in the parameter-sensitivity charts concerns Olsson–Glass. In the local sensitivity comparisons, Olsson–Glass often appears weaker at the baseline setting than at the two tested side settings. In other words, the middle baseline bar is lower than the bars corresponding to the altered parameter values. This pattern should be interpreted cautiously, because the parameter-sensitivity analysis is local and descriptive rather than factorial. Nevertheless, it suggests that Olsson–Glass performance may be sensitive to parameter configuration and that the baseline setting may not be the most favorable condition for this rule. This is important because the main survival results show that Olsson–Glass is highly competitive in the coherence-dominant regime, especially at coin bias 0.7, but the local sensitivity plots suggest that its relative performance may depend on the exact parameter environment.

The parameter-sensitivity results also reinforce the intermediate status of Hartmann–

Trpin. Although Hartmann–Trpin never survived to the final generation, it consistently persisted longer than Shogenji and Olsson–Glass in the Bonus-dominant regime. This was not only visible in the overall extinction summary, but also remained visible across the parameter-sensitivity averages. Thus, Hartmann–Trpin should not be interpreted as simply unsuccessful. It was unsuccessful in the sense that it never became a final survivor, but it was comparatively persistent among the losing groups. Its failure was therefore a failure of final dominance, not a failure of short-term or medium-term competitiveness.

4.4.1 Effect of Sequence Length

Sequence length was tested at 150, 200, and 250 observations, with threshold fixed at 0.90 and elitism fixed at 0.05. The baseline was sequence length 200.

Changing sequence length did not change final survivor identity. Bonus agents continued to dominate the Bonus-dominant regime, and coherence agents continued to dominate the coherence-dominant regime. The only visible effect was in extinction timing.

In the Bonus-dominant regime, the latest average extinction among losing groups was 18 generations at the baseline sequence length of 200. It was 17 generations at sequence length 150 and 16 generations at sequence length 250. These differences are small. They suggest that, within this tested range, sequence length had only a limited effect on how quickly Bonus agents achieved full dominance.

In the coherence-dominant regime, the latest average extinction was 17 generations for all three sequence lengths. This suggests that, in this regime, changing the available sequence length from 150 to 250 did not visibly change the timing of dominance.

To make the two survival regimes easier to visualize, Figure 4.5 presents schematic reconstructions of average population development in the Bonus-dominant and coherence-dominant regimes.

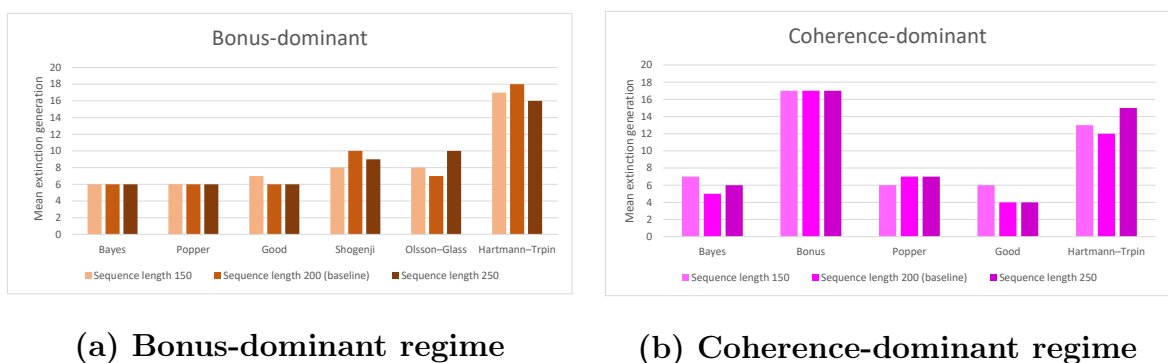


Figure 4.5: Schematic reconstructions of average population development under sequence-length variation.

These plots should not be interpreted as direct generation-by-generation records from individual simulations. The actual simulation trajectories were non-linear and often contained fluctuations in population size before final dominance stabilized. Instead, the figures summarize the average extinction structure: extinct groups are shown declining from their initial population of 50 agents to zero by their mean extinction generation, while surviving groups fill the remaining population. The purpose of the plots is therefore illustrative. They show the approximate timing of extinction and dominance implied by the summary tables, not the exact evolutionary path of any single run.

Overall, the sequence-length analysis supports the robustness of the main result. The final survival pattern did not depend on whether agents were evaluated with 150, 200, or 250 observations. Any effect of sequence length was limited to small changes in extinction timing.

4.4.2 Effect of Threshold

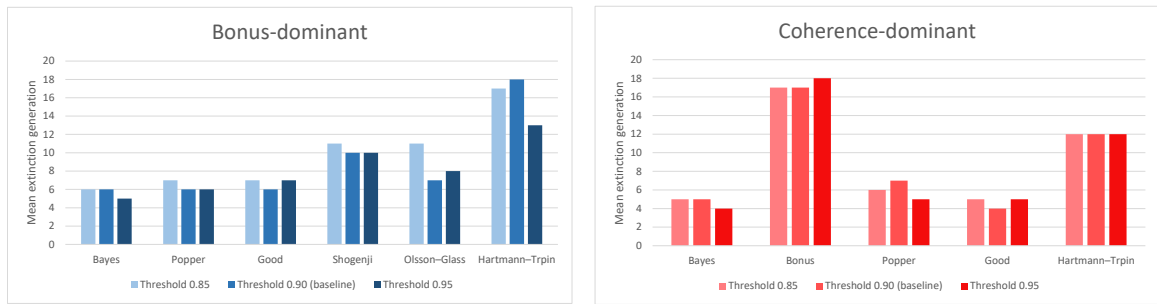
Threshold was tested at 0.85, 0.90, and 0.95, with sequence length fixed at 200 and elitism fixed at 0.05. The baseline was threshold 0.90.

As with sequence length, changing the threshold did not change final survivor identity. The main pattern remained the same: Bonus agents dominated in the Bonus-dominant regime, while coherence agents dominated in the coherence-dominant regime.

The timing effects were somewhat more visible than in the sequence-length analysis. In the Bonus-dominant regime, the latest average extinction was 18 generations at the baseline threshold of 0.90. It was 17 generations at threshold 0.85 and 13 generations at threshold 0.95. This suggests that, descriptively, the stricter threshold was associated with faster final elimination of losing groups in the Bonus-dominant regime.

In the coherence-dominant regime, the latest average extinction remained very stable. It was 17 generations at thresholds 0.85 and 0.90, and 18 generations at threshold 0.95. Thus, threshold variation had little effect on coherence-dominant extinction timing.

To make the two survival regimes easier to visualize, Figure 4.6 presents schematic reconstructions of average population development in the Bonus-dominant and coherence-dominant regimes.



(a) Bonus-dominant regime

(b) Coherence-dominant regime

Figure 4.6: Threshold sensitivity by survival regime. Panel (a) shows extinction timing in the Bonus-dominant regime, and panel (b) shows extinction timing in the coherence-dominant regime.

The threshold analysis therefore supports two conclusions. First, the main survival pattern is robust across the tested confidence thresholds. Second, threshold may affect extinction timing in some regimes, especially the Bonus-dominant regime, but the evidence is descriptive rather than conclusive.

4.4.3 Effect of Elitism

Elitism was tested at 0.00, 0.05, and 0.20, with sequence length fixed at 200 and threshold fixed at 0.90. The baseline was elitism 0.05.

Elitism produced the clearest timing differences among the parameter tests. In the Bonus-dominant regime, the latest average extinction was 18 generations at baseline elitism 0.05. Without elitism, the latest average extinction increased to 21 generations. Under stronger elitism, 0.20, it decreased to 13 generations. This suggests that stronger elitism accelerated dominance by the surviving group, while no elitism allowed losing groups to persist longer.

The same broad pattern appeared in the coherence-dominant regime. The latest average extinction was 17 generations at both no elitism and baseline elitism, but decreased to 12 generations under strong elitism. Thus, stronger elitism appears to have accelerated extinction timing in both regimes.

To make the two survival regimes easier to visualize, Figure 4.7 presents schematic reconstructions of average population development in the Bonus-dominant and coherence-dominant regimes.

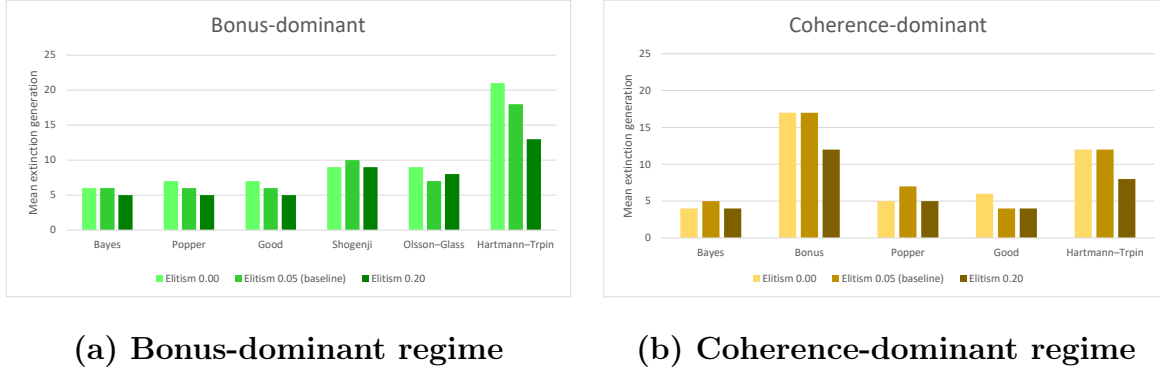


Figure 4.7: Elitism sensitivity by survival regime. Panel (a) shows extinction timing in the Bonus-dominant regime, and panel (b) shows extinction timing in the coherence-dominant regime.

This result is consistent with the role of elitism in the evolutionary algorithm. Elitism preserves top-performing agents and therefore can intensify selection. In this simulation, stronger elitism did not change which group survived, but it did appear to affect how quickly losing groups disappeared. Across the parameter-sensitivity analyses, the most important result is that parameter changes did not reverse the main survival regimes. However, the secondary timing patterns are still informative. Hartmann–Trpin repeatedly remained more persistent than Shogenji and Olsson–Glass when those coherence agents were in unfavorable Bonus-dominant conditions. At the same time, Olsson–Glass showed signs of parameter sensitivity, since its performance was often weaker at baseline than under the adjacent tested parameter settings. These observations should be treated as descriptive tendencies, but they indicate that extinction timing contains information that is not visible from final survival alone.

4.5 Special Case: Fully Biased Coin

The final result concerns the fully biased coin condition, where the true coin bias was 0.0. This condition belongs to the Bonus-dominant regime, but it appears to differ from the other Bonus-dominant biases in how quickly Bonus achieved full dominance.

Table 4.5 extinction timing for the fully biased coin condition with the other Bonus-dominant conditions. The comparison includes all seven parameter configurations. The final column reports the average generation at which Bonus achieved full dominance, calculated as the latest extinction generation among all non-Bonus groups in each simulation.

Table 4.5: Extinction timing in the fully biased coin condition compared with other Bonus-dominant biases.

Condition	N	Bayes	Popper	Good	Shogenji	Olsson–Glass	Hartmann–Trpin	Mean losing-group extinction	Bonus full dominance
Fully biased coin 0.0	7	2	3	3	5	5	5	4	5
Other Bonus-dominant biases	42	6	7	7	10	9	18	10	19
Difference	–	4	4	4	5	4	13	6	14

The comparison in Table 4.5 shows that the fully biased condition produced much faster extinction of competing groups. In the fully biased condition, the average extinction generations of losing groups were very low: Bayes disappeared around generation 2, Popper and Good around generation 3, and Shogenji, Olsson–Glass, and Hartmann–Trpin around generation 5. The mean extinction across losing groups was 4 generations. Bonus achieved full dominance around generation 5.

By contrast, in the other Bonus-dominant biases, losing groups persisted longer. Hartmann–Trpin in particular survived much longer in these conditions, with a mean extinction generation of 18. The average generation of full Bonus dominance was 19, compared with only 5 in the fully biased condition. This contrast is illustrated in Figure 4.8.

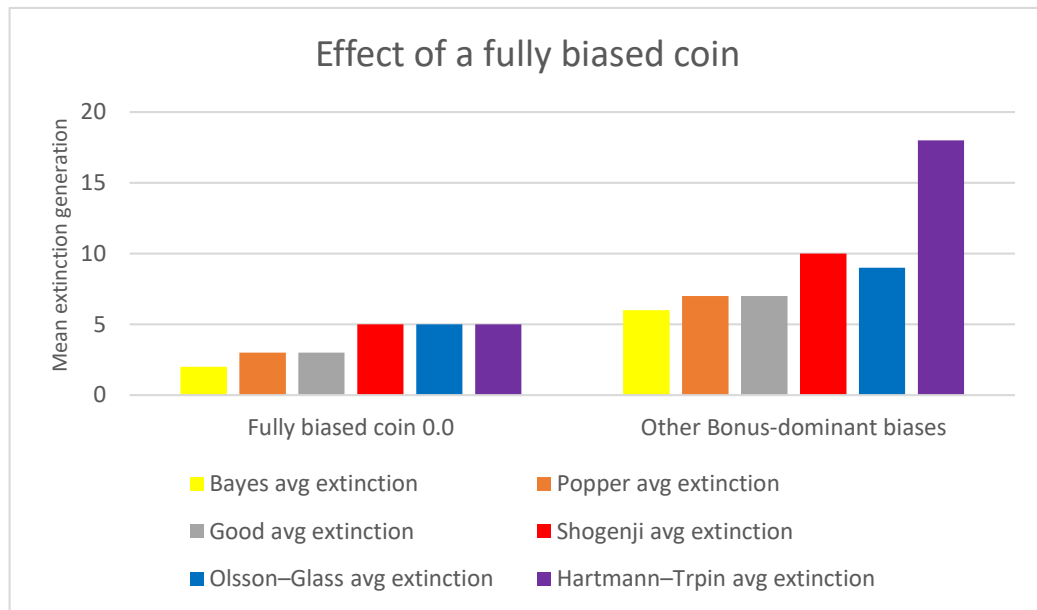


Figure 4.8: Comparison of the fully biased coin condition with other Bonus-dominant biases. The figure shows the much earlier extinction of losing groups and the earlier achievement of full Bonus dominance in the fully biased coin condition.

This result suggests that the fully biased coin condition is easier for Bonus agents to dominate. When the evidence is maximally one-sided, the Bonus rule appears to identify and reinforce the correct direction very quickly. However, this interpretation should remain descriptive. The fully biased comparison is based on seven simulations,

one for each parameter configuration. Therefore, it should be treated as an observed pattern in the simulation rather than as a statistically established effect.

4.6 Summary of Empirical Findings

The results support four main findings.

First, the simulations produced a strong survival pattern across all 70 runs. Bonus agents dominated at coin biases 0.0, 0.1, 0.2, 0.4, 0.6, 0.8, and 0.9. Coherence-based agents dominated at coin biases 0.3, 0.5, and 0.7.

Second, Shogenji and Olsson–Glass agents should be interpreted as the surviving coherence category in the coherence-dominant regime. Shogenji survived in all coherence-dominant simulations, while Olsson–Glass survived in most and remained highly competitive even when it became extinct.

Third, Bayes, Popper, Good, and Hartmann–Trpin agents did not survive to the final generation in any condition. However, Hartmann–Trpin differed from Bayes, Popper, and Good because it often persisted longer before extinction. Fourth, parameter changes did not substantially alter the final survival pattern. Sequence length, threshold, and elitism mainly affected extinction timing. Among these, elitism produced the clearest timing differences: stronger elitism was associated with earlier extinction of losing groups. Overall, the results show that the evolutionary model generated a robust division between Bonus-dominant and coherence-dominant regimes. The main determinant of survivor identity was not the tested parameter setting, but the true coin-bias condition.

Chapter 5

General Discussion

5.1 Interpretation of the Main Pattern

The main result of the study is the emergence of two survival regimes. Bonus agents dominated under most coin-bias conditions, while coherence-based agents dominated under coin biases 0.3, 0.5, and 0.7. This suggests that different update rules may be advantageous under different evidence environments.

The Bonus rule performed especially well when the evidence stream was strongly informative in a direction that could be captured by the running empirical frequency. In these conditions, the rule could quickly reinforce hypotheses close to the observed pattern. This may explain why Bonus agents dominated under most true-bias values, especially in the fully biased condition 0.0, where the evidence was maximally one-sided.

The coherence-based rules performed best under the intermediate and unbiased conditions. This is theoretically interesting because these conditions are less straightforward. In particular, the unbiased condition 0.5 does not provide a strong directional signal toward one extreme hypothesis. Under such conditions, coherence-sensitive updating may have an advantage because it evaluates the structure of the accumulated evidence pattern rather than only rewarding immediate empirical closeness.

The near balance between Shogenji and Olsson–Glass at coin bias 0.5 is also noteworthy. It suggests that under maximally ambiguous evidence, neither coherence rule clearly eliminates the other. Instead, both remain viable. At 0.3 and 0.7, the balance shifts, with Shogenji stronger at 0.3 and Olsson–Glass stronger at 0.7. This asymmetry should not be overinterpreted, but it indicates that the two coherence measures do not behave identically under evolutionary selection.

The failure of standard Bayes to survive does not mean that Bayesian updating is irrational or generally weak. In the simulation, all rules share a Bayesian core, and the non-standard rules modify this core through explanatory or coherence-sensitive

mechanisms. The result therefore suggests that, under the specific fitness function used here, additional explanatory or coherence-guided mechanisms can outperform pure conditionalization in terms of evolutionary survival. This should be understood as a result about the implemented environment, not as a general rejection of Bayesian reasoning.

5.2 Relation to Explanation and Coherence

The findings are consistent with the broader motivation of the thesis: belief-update rules can be evaluated dynamically by asking how they behave under repeated evidence and selection. The results suggest that explanatory and coherence-sensitive modifications can matter, but their success is conditional.

The empirical Bonus rule performed very strongly in most conditions. This supports the idea that an explanatory or frequency-sensitive preference can be powerful when the observed evidence pattern provides a clear guide to the true hypothesis. The Bonus rule does not merely update by likelihood. It gives additional support to hypotheses closest to the observed frequency, and this appears to be highly effective across most tested coin-bias values.

The coherence rules performed selectively. They did not dominate everywhere, but they were successful in a specific subset of conditions. This is important because coherence is often discussed as an epistemic virtue, but the simulation suggests that its usefulness may depend on the evidential environment. In the present model, coherence-based updating was strongest under intermediate or ambiguous bias conditions, not under all conditions.

Hartmann–Trpin agents are especially interesting in this context. They never survived to the final generation, but they frequently persisted longer than other losing groups. This suggests that their coherence criterion may have provided some resistance to immediate extinction without producing enough selective advantage to dominate. In future work, it would be useful to examine whether different implementations of the Hartmann–Trpin rule, especially different bonus sizes or posterior-based variants, would change this outcome.

Overall, the results support a conditional view of epistemic update rules. No rule should be assumed to be universally superior. Instead, the performance of a rule depends on the structure of the evidence, the fitness criterion, and the selection environment.

5.3 Methodological Limitations

The most important limitation of the study concerns the parameter analysis. The thesis uses a local sensitivity analysis rather than a full factorial parameter analysis. Sequence length, threshold, and elitism were each varied one at a time around a baseline configuration. This design made the simulations computationally manageable, but it means that interaction effects between parameters were not tested.

For example, the study does not examine every possible combination of sequence length, threshold, and elitism. It therefore cannot answer whether a high threshold combined with strong elitism and a long sequence length would produce different dynamics from those observed in the one-factor-at-a-time design. A full factorial design would provide a stronger test of parameter interactions, but it would require many more simulations. Given computational and time limitations, the local sensitivity design was a practical compromise. This limitation is especially relevant for interpreting the comparison between Shogenji and Olsson–Glass. In the overall results, Shogenji has the higher average final population across the coherence-dominant conditions. However, some of the parameter-sensitivity charts suggest that Olsson–Glass may perform better away from the baseline setting than it does at the baseline itself. Because the analysis varies only one parameter at a time, it cannot determine whether Olsson–Glass would perform better under combined non-baseline settings, such as a different threshold together with a different elitism level. A full factorial design could therefore change the apparent balance between Shogenji and Olsson–Glass. A related limitation concerns the asymmetry between the 0.3 and 0.7 bias conditions. Since these conditions are mirror cases, one might expect the relative performance of Shogenji and Olsson–Glass to be symmetric as well. The observed shift from Shogenji dominance at 0.3 to Olsson–Glass dominance at 0.7 should therefore be treated as tentative until tested with a larger number of repeated runs. For this reason, the thesis should not conclude that Shogenji is generally superior to Olsson–Glass. The safer conclusion is that Shogenji performed better on average within the tested local design, while Olsson–Glass remained highly competitive and may be more sensitive to parameter configuration.

A second limitation is that the parameter-specific analyses have small subgroup sizes. Once the simulations are divided by survival regime and parameter value, the number of cases becomes limited. This is especially true for the coherence-dominant regime, which contains only three coin-bias values. Therefore, the parameter analyses should be treated as descriptive indicators of timing tendencies, not as statistically conclusive evidence.

A third limitation concerns the selection-only evolutionary design. The model intentionally excludes mutation and crossover. This was methodologically appropriate because the goal was to compare fixed update-rule groups, not to generate new hybrid

rules. However, it also limits the evolutionary interpretation of the model. In a more open-ended evolutionary algorithm, mutation and recombination could create new update strategies. Such a model might reveal whether hybrid rules combining features of Bonus and coherence-based updating could outperform the fixed rules used here.

A fourth limitation concerns the evidence environment. The simulation uses repeated coin-toss evidence and a discrete hypothesis space of possible coin biases. This environment is useful because it is controlled, transparent, and mathematically tractable. However, it is also simpler than many real epistemic environments. Scientific reasoning often involves multi-dimensional hypotheses, dependent evidence, background assumptions, noisy measurements, and structured causal relations. The results therefore show how the update rules behave in a controlled binary setting, not how they would necessarily behave in all scientific contexts.

A fifth limitation concerns the fitness function. Fitness combines time-to-threshold with a Brier-score penalty. This is defensible because it captures both speed and accuracy, but it is still one possible operationalization of epistemic performance. Different threshold values, different penalty weights, or different scoring rules might produce different outcomes. The study partially addresses this by testing several thresholds, but future work could examine a wider range of scoring and decision criteria.

Finally, the generation-by-generation visualizations should be interpreted with care. The original simulation plots show non-linear population dynamics, including fluctuations before final dominance. Schematic area plots based on mean extinction timing can be useful for illustration, but they are not direct records of the actual population trajectories. For transparency, the full set of original simulation plots should be included in the electronic appendix.

5.4 Future Work

Future work could extend the present study in several directions. First, a full factorial parameter design would allow the effects of sequence length, threshold, and elitism to be tested more systematically. This would make it possible to determine whether combinations of parameters interact in ways that are not visible in the local sensitivity analysis.

Second, more repeated simulations could be run for each parameter-bias combination. This would allow stronger statistical analysis of stochastic variation. In the current study, the results are descriptive. With more replications, future work could estimate confidence intervals, compare distributions of extinction times, and test whether observed differences are stable across repeated runs.

Third, the model could be extended beyond a binary coin-toss environment. More

complex evidence environments could include multi-valued evidence, dependent observations, noisy signals, or causal structures. This would make it possible to test whether coherence-based advantages persist in richer epistemic settings.

Fourth, future versions of the evolutionary algorithm could introduce controlled mutation or recombination. This would change the research question from comparing fixed rule groups to discovering new hybrid update strategies. Such an extension would need to preserve interpretability, but it could reveal whether combinations of explanatory and coherence-sensitive mechanisms outperform the fixed rules tested here.

Fifth, the implementation of coherence rules could be expanded. In the current model, Shogenji, Olsson–Glass, and Hartmann–Trpin measures are used as selection devices for assigning a fixed bonus. Future work could test alternative implementations, such as variable bonus sizes, posterior-based coherence scores, or continuous coherence weighting. This would help determine whether the observed performance of the coherence agents depends on the measures themselves or on the way they were operationalized. A further extension would be to run a coherence-only evolutionary experiment. In the present design, Shogenji, Olsson–Glass, and Hartmann–Trpin agents compete not only against one another, but also against Bayesian, explanatory, and empirical-bonus agents. Hartmann–Trpin never survives in the full competition, but it often persists longer than Shogenji and Olsson–Glass when the coherence agents are in unfavorable conditions. This raises the possibility that Hartmann–Trpin may be more competitive in a restricted population containing only coherence-based agents. Such an experiment would help distinguish two possibilities: whether Hartmann–Trpin is genuinely weaker as a coherence rule, or whether it is disadvantaged mainly by the presence of stronger non-coherence competitor in the full evolutionary environment.

5.5 Concluding Interpretation

The results show that evolutionary survival of belief-update rules depends strongly on the evidential environment. The Bonus rule dominated under most true-bias values, especially when the evidence was strongly directional. Coherence-based rules dominated under the intermediate and unbiased conditions 0.3, 0.5, and 0.7. This suggests that coherence-sensitive updating may be especially useful when evidence is less straightforward and when the structure of the accumulated evidence matters.

At the same time, the study does not show that any single rule is universally best. The performance of each rule depends on the interaction between the evidence environment, the fitness function, and the evolutionary selection mechanism. The main contribution of the thesis is therefore not the identification of one universally superior update rule, but the demonstration that Bayesian, explanatory, and coherence-sensitive

rules can be placed into a shared computational selection framework and compared through survival, extinction, and dominance dynamics.

Conclusion

This thesis examined belief revision under uncertainty by comparing several epistemic update rules in a controlled computational environment. The central question was not only which rule can eventually identify the true hypothesis, but also which rule can do so quickly and accurately enough to be useful. In this sense, the thesis treated belief revision as a dynamic process rather than as a single final judgment. Agents repeatedly received evidence, updated their beliefs about the true bias of a coin, and were evaluated through an evolutionary selection mechanism based on speed and accuracy. This made it possible to compare Bayesian, explanatory, and coherence-sensitive approaches within one shared framework.

The main contribution of the thesis lies in transforming philosophical and epistemological ideas into computationally testable rules. Coherence and explanation are often discussed as abstract epistemic virtues, especially in philosophy of science and formal epistemology. However, this thesis showed that they can also be operationalized in a simulation model. Shogenji's measure, the Glass–Olsson measure, and the Hartmann–Trpin measure were not treated only as theoretical formulas, but as tools for guiding belief revision. Similarly, explanatory reasoning was not discussed only as a philosophical alternative to Bayesianism, but was implemented through bonus and weighting mechanisms. In this way, the thesis connected normative questions about rational belief with empirical questions about performance, survival, and adaptation.

The results show that no single belief-update rule should be understood as universally superior. The empirical Bonus rule dominated under most tested true-bias conditions, especially where the evidence was strongly directional and the observed frequency provided a clear guide to the true hypothesis. Coherence-based rules, especially Shogenji and Glass–Olsson, dominated under the intermediate and unbiased conditions 0.3, 0.5, and 0.7. This suggests that coherence-sensitive updating may be most useful when the evidence is less straightforward and when the structure of the accumulated evidence matters. Hartmann–Trpin agents did not survive to the final generation, but their persistence under some unfavorable conditions suggests that their role should not be dismissed too quickly. Their failure was a failure of final dominance, not necessarily a complete failure of epistemic usefulness.

The results also clarify the role of Bayesian updating in the thesis. The fact that

standard Bayesian agents did not survive in the final populations does not imply that Bayesian reasoning is irrational or weak in general. On the contrary, Bayesian updating remained the common core of the implemented rules. The explanatory and coherence-sensitive rules modified this Bayesian core by adding additional mechanisms. The results therefore suggest that, in the specific environment modeled here, such additional mechanisms can provide selective advantages under certain conditions. This conclusion should be interpreted carefully: the thesis does not reject Bayesianism, but it shows that Bayesian conditionalization can be fruitfully compared with rules that incorporate explanatory or coherence-based considerations.

The interdisciplinary character of the thesis is essential to this contribution. From philosophy and epistemology, the thesis takes the concepts of coherence, explanation, justification, and rational belief revision. From mathematics, it uses probability theory, Bayesian updating, coherence measures, explanatory measures, normalization, and scoring rules. From computer science, it draws on simulation, algorithmic implementation, and evolutionary computation. From cognitive science and psychology, it connects to the study of reasoning under uncertainty and the trade-off between fast and accurate judgment. Finally, from biology, it borrows the logic of evolutionary selection, where rule groups survive, decline, or disappear according to their performance across generations. The thesis therefore does not belong fully to only one discipline. Its value comes from placing these different perspectives into contact with one another.

This interdisciplinary approach is important because real reasoning under uncertainty is itself not purely philosophical, mathematical, psychological, or computational. Scientific and everyday reasoning involve formal structure, cognitive limitations, explanatory preferences, and adaptive pressure at the same time. A rule may be mathematically elegant but too slow to be useful. Another rule may be fast but insufficiently accurate. A third rule may work well only in certain environments. By evaluating update rules through both speed and accuracy, this thesis reflects the practical demands of reasoning in situations where evidence is incomplete, noisy, and time-sensitive.

At the same time, the study has clear limitations. The simulation used a simplified coin-toss environment, and the parameter analysis was local rather than fully factorial. The results are descriptive and should not be interpreted as final statistical proof that one family of rules is generally better than another. The model also compared fixed rule groups without mutation, recombination, communication, or social learning. These simplifications were useful because they made the experiment controlled and interpretable, but they also limit the generality of the conclusions.

Future research could extend the model in several directions. A full factorial parameter design could test interactions between sequence length, threshold, and elitism more systematically. More repetitions could support stronger statistical analysis. More complex evidence environments could test whether coherence-sensitive advantages per-

sist beyond simple binary evidence. The evolutionary model could also be expanded to include mutation, recombination, or hybrid update rules. Finally, coherence measures could be implemented in alternative ways, for example through variable bonus sizes, posterior-based coherence scores, or continuous coherence weighting. Such extensions would help determine whether the observed results depend mainly on the measures themselves, on the implementation choices, or on the structure of the simulated environment.

In conclusion, this thesis shows that belief revision can be studied as an interdisciplinary problem that connects philosophy, mathematics, computer science, and evolutionary modeling. It demonstrates that epistemic rules can be evaluated not only by abstract argument, but also by their behavior in simulated environments. The central lesson is conditional rather than absolute: different rules succeed under different evidential conditions. Coherence and explanation are not universal replacements for Bayesian updating, but they can function as meaningful additions to probabilistic belief revision. By placing these rules into a shared evolutionary framework, the thesis contributes to a broader understanding of how agents may track truth under uncertainty, and how rational inquiry can be evaluated as a balance between accuracy, speed, and adaptation.

Bibliography

- Adami, C., Schossau, J., & Hintze, A. (2016). Evolutionary game theory using agent-based methods. *Physics of Life Reviews*, 19, 1–26. <https://doi.org/10.1016/j.plrev.2016.08.015>
- Arcuri, A., & Briand, L. (2014). A hitchhiker’s guide to statistical tests for assessing randomized algorithms in software engineering. *Software Testing, Verification and Reliability*, 24(3), 219–250. <https://doi.org/10.1002/stvr.1486>
- Baccini, E., Christoff, Z., van Leeuwen, L., & Verbrugge, R. (2026). Tracking the truth by selecting good data: Coherence measures and data selection. *Synthese*, 207(3), Article 88. <https://doi.org/10.1007/s11229-026-05453-9>
- Bayes, T. (1763). An essay towards solving a problem in the doctrine of chances. *Philosophical Transactions of the Royal Society of London*, 53, 370–418. <https://doi.org/10.1098/rstl.1763.0053>
- Blanshard, B. (1939). *The Nature of Thought* (Vols. 2). George Allen & Unwin.
- Blickle, T., & Thiele, L. (1996). A comparison of selection schemes used in evolutionary algorithms. *Evolutionary Computation*, 4(4), 361–394. <https://doi.org/10.1162/evco.1996.4.4.361>
- BonJour, L. (1985). *The Structure of Empirical Knowledge*. Harvard University Press.
- Bradley, F. H. (1893). *Appearance and Reality: A Metaphysical Essay*. Swan Sonnenschein & Co.
- Brier, G. W. (1950). Verification of forecasts expressed in terms of probability. *Monthly Weather Review*, 78(1), 1–3. [https://doi.org/10.1175/1520-0493\(1950\)078<0001:VOFEIT>2.0.CO;2](https://doi.org/10.1175/1520-0493(1950)078<0001:VOFEIT>2.0.CO;2)
- Bromberger, S. (1966). Why-questions. In R. G. Colodny (Ed.), *Mind and cosmos: Essays in contemporary science and philosophy* (pp. 86–111). University of Pittsburgh Press.
- Czitrom, V. (1999). One-factor-at-a-time versus designed experiments. *The American Statistician*, 53(2), 126–131. <https://doi.org/10.1080/00031305.1999.10474445>
- De Luca, G., Suryapranata, H., Ottervanger, J. P., & Antman, E. M. (2004). Time delay to treatment and mortality in primary angioplasty for acute myocardial infarction: Every minute of delay counts. *Circulation*, 109(10), 1223–1225. <https://doi.org/10.1161/01.CIR.0000121424.76486.20>

- Douven, I. (2013). Inference to the best explanation, Dutch books, and inaccuracy minimisation. *The Philosophical Quarterly*, 63(252), 428–444. <https://doi.org/10.1111/1467-9213.12032>
- Douven, I. (2019). Optimizing group learning: An evolutionary computing approach. *Artificial Intelligence*, 275, 235–251. <https://doi.org/10.1016/j.artint.2019.06.002>
- Douven, I. (2025). Adaptive updating: Ecological rationality meets reinforcement learning. *Minds and Machines*, 35(3), 1–23. <https://doi.org/10.1007/s11023-025-09735-y>
- Douven, I., & Schupbach, J. N. (2015). Probabilistic alternatives to Bayesianism: The case of explanationism. *Frontiers in Psychology*, 6, Article 459. <https://doi.org/10.3389/fpsyg.2015.00459>
- Douven, I., & Wenmackers, S. (2017). Inference to the best explanation versus Bayes’s rule in a social setting. *The British Journal for the Philosophy of Science*, 68(2), 535–570. <https://doi.org/10.1093/bjps/axv025>
- Douven, I. (1999). Inference to the best explanation made coherent. *Philosophy of Science*, 66(S3), S424–S435. <https://doi.org/10.1086/392743>
- Douven, I. (2016). Explanation, updating, and accuracy. *Journal of Cognitive Psychology*, 28(8), 1004–1012. <https://doi.org/10.1080/20445911.2016.1230122>
- Drugowitsch, J., DeAngelis, G. C., Angelaki, D. E., & Pouget, A. (2015). Tuning the speed–accuracy trade-off to maximize reward rate in multisensory decision-making. *eLife*, 4, Article e06678. <https://doi.org/10.7554/eLife.06678>
- Drugowitsch, J., Mendonça, A. G., Mainen, Z. F., & Pouget, A. (2019). Learning optimal decisions with confidence. *Proceedings of the National Academy of Sciences of the United States of America*, 116(49), 24872–24880. <https://doi.org/10.1073/pnas.1906787116>
- Eiben, A. E., & Smith, J. E. (2015). *Introduction to Evolutionary Computing* (2nd ed.). Springer. <https://doi.org/10.1007/978-3-662-44874-8>
- Glass, D. H. (2002). Coherence, Explanation and Bayesian Networks. In M. O’Neill, R. F. E. Sutcliffe, C. Ryan, M. Eaton, & N. J. L. Griffith (Eds.), *Artificial Intelligence and Cognitive Science* (pp. 177–182, Vol. 2464). Springer. https://doi.org/10.1007/3-540-45750-X_23
- Glass, D. H. (2007). Coherence measures and inference to the best explanation. *Synthese*, 157(3), 275–296. <https://doi.org/10.1007/s11229-006-9055-7>
- Gneiting, T., & Raftery, A. E. (2007). Strictly proper scoring rules, prediction, and estimation. *Journal of the American Statistical Association*, 102(477), 359–378. <https://doi.org/10.1198/0162145060000001437>
- Goldberg, D. E., & Deb, K. (1991). A Comparative Analysis of Selection Schemes Used in Genetic Algorithms. In G. J. E. Rawlins (Ed.), *Foundations of Genetic*

- Algorithms* (pp. 69–93, Vol. 1). Morgan Kaufmann. <https://doi.org/10.1016/B978-0-08-050684-5.50008-2>
- Good, I. J. (1960). Weight of evidence, corroboration, explanatory power, information and the utility of experiments. *Journal of the Royal Statistical Society: Series B (Methodological)*, 22(2), 319–331. <https://doi.org/10.1111/j.2517-6161.1960.tb00378.x>
- Harman, G. H. (1965). The inference to the best explanation. *The Philosophical Review*, 74(1), 88–95. <https://doi.org/10.2307/2183532>
- Harris, A. J. L., & Hahn, U. (2009). Bayesian rationality in evaluating multiple testimonies: Incorporating the role of coherence. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 35(5), 1366–1373. <https://doi.org/10.1037/a0016567>
- Hartmann, S., & Trpin, B. (2023). Coherence of information: What it is and why it matters. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 45, 3617–3623. <https://escholarship.org/uc/item/4w7343x9>
- Hartmann, S., & Trpin, B. (2024). A new posterior probability-based measure of coherence. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46, 1494–1500. <https://escholarship.org/uc/item/30p8x5xh>
- Hartmann, S., & Trpin, B. (2026a). Why coherence matters [Advance online publication]. *The Journal of Philosophy*. <https://doi.org/10.5840/jphil20264117>
- Hartmann, S., & Trpin, B. (2026b). Coherence-based measures of explanatory power. *Philosophical Studies*, 183, 1817–1843. <https://doi.org/10.1007/s11098-026-02517-x>
- Hempel, C. G., & Oppenheim, P. (1948). Studies in the logic of explanation. *Philosophy of Science*, 15(2), 135–175. <https://doi.org/10.1086/286983>
- Hempel, C. G. (1945). Studies in the logic of confirmation (i.) *Mind*, 54(213), 1–26. <https://doi.org/10.1093/mind/liv.213.1>
- Holland, J. H. (1975). *Adaptation in Natural and Artificial Systems: An Introductory Analysis with Applications to Biology, Control, and Artificial Intelligence*. University of Michigan Press.
- Hsiang, S., Allen, D., Annan-Phan, S., Bell, K., Bolliger, I., Chong, T., Druckenmiller, H., Huang, L. Y., Hultgren, A., Krasovich, E., Lau, P., Lee, J., Rolf, E., Tseng, J., & Wu, T. (2020). The effect of large-scale anti-contagion policies on the COVID-19 pandemic. *Nature*, 584(7820), 262–267. <https://doi.org/10.1038/s41586-020-2404-8>
- Jaynes, E. T. (2003). *Probability theory: The logic of science*. Cambridge University Press.
- Joachim, H. H. (1906). *The Nature of Truth: An Essay*. Clarendon Press.

- Kolmogorov, A. N. (1950). *Foundations of the theory of probability* [Translated by N. Morrison. Original work published 1933]. Chelsea Publishing Company.
- Koscholke, J., Schippers, M., & Stegmann, A. (2019). New hope for relative overlap measures of coherence. *Mind*, 128(512), 1261–1284. <https://doi.org/10.1093/mind/fzy037>
- Laplace, P.-S. (1986). Memoir on the probability of the causes of events (S. M. Stigler, Trans.). *Statistical Science*, 1(3), 364–378. <https://doi.org/10.1214/ss/1177013621> (Original work published 1774)
- Lehrer, K. (1990). *Theory of Knowledge*. Westview Press.
- Lewis, C. I. (1946). *An Analysis of Knowledge and Valuation*. Open Court.
- Liesefeld, H. R., & Janczyk, M. (2019). Combining speed and accuracy to control for speed–accuracy trade-offs(?) *Behavior Research Methods*, 51(1), 40–60. <https://doi.org/10.3758/s13428-018-1076-x>
- Miller, B. L., & Goldberg, D. E. (1995). Genetic algorithms, tournament selection, and the effects of noise. *Complex Systems*, 9(3), 193–212.
- Mills, K. L., Filliben, J. J., & Haines, A. L. (2015). Determining relative importance and effective settings for genetic algorithm control parameters. *Evolutionary Computation*, 23(2), 309–342. https://doi.org/10.1162/EVCO_a_00137
- Nowak, M. A., & Sigmund, K. (2004). Evolutionary dynamics of biological games. *Science*, 303(5659), 793–799. <https://doi.org/10.1126/science.1093411>
- Olsson, E. J. (2002). What is the problem of coherence and truth? *The Journal of Philosophy*, 99(5), 246–272. <https://doi.org/10.2307/3655648>
- Olsson, E. J. (2005). *Against Coherence: Truth, Probability, and Justification*. Oxford University Press. <https://doi.org/10.1093/0199279993.001.0001>
- Pei, S., Kandula, S., & Shaman, J. (2020). Differential effects of intervention timing on COVID-19 spread in the United States. *Science Advances*, 6(49), Article eabd6370. <https://doi.org/10.1126/sciadv.abd6370>
- Peirce, C. S. (1931–1958). *Collected Papers of Charles Sanders Peirce* (C. Hartshorne, P. Weiss, & A. W. Burks, Eds.; Vols. 8). Harvard University Press.
- Popper, K. R. (1959). *The logic of scientific discovery* (2nd ed.) [Original work published 1934]. Routledge; Kegan Paul.
- Quine, W. V. O., & Ullian, J. S. (1970). *The Web of Belief*. Random House.
- Ratcliff, R., & McKoon, G. (2008). The diffusion decision model: Theory and data for two-choice decision tasks. *Neural Computation*, 20(4), 873–922. <https://doi.org/10.1162/neco.2008.12-06-420>
- Rescher, N. (1973). *The Coherence Theory of Truth*. Clarendon Press.
- Russell, B. (1907). On the nature of truth. *Proceedings of the Aristotelian Society*, 7(1), 28–49. <https://doi.org/10.1093/aristotelian/7.1.28>

- Schupbach, J. N., & Sprenger, J. (2011). The logic of explanatory power. *Philosophy of Science*, 78(1), 105–127. <https://doi.org/10.1086/658111>
- Seymour, C. W., Gesten, F., Prescott, H. C., Friedrich, M. E., Iwashyna, T. J., Phillips, G. S., Lemeshow, S., Osborn, T., Terry, K. M., & Levy, M. M. (2017). Time to treatment and mortality during mandated emergency care for sepsis. *The New England Journal of Medicine*, 376(23), 2235–2244. <https://doi.org/10.1056/NEJMoA1703058>
- Shogenji, T. (1999). Is coherence truth conducive? *Analysis*, 59(4), 338–345. <https://doi.org/10.1093/analys/59.4.338>
- ten Broeke, G., van Voorn, G., & Ligtenberg, A. (2016). Which sensitivity analysis method should i use for my agent-based model? *Journal of Artificial Societies and Social Simulation*, 19(1), Article 5. <https://doi.org/10.18564/jasss.2857>
- Thagard, P. (1989). Explanatory coherence. *Behavioral and Brain Sciences*, 12(3), 435–502. <https://doi.org/10.1017/S0140525X00057319>
- Trpin, B., & Justin, M. (2025). Coherence as a constraint on scientific inquiry. *Synthese*, 206(4), Article 200. <https://doi.org/10.1007/s11229-025-05281-3>
- van Fraassen, B. C. (1989). *Laws and symmetry*. Oxford University Press.
- Wald, A. (1945). Sequential tests of statistical hypotheses. *The Annals of Mathematical Statistics*, 16(2), 117–186. <https://doi.org/10.1214/aoms/1177731118>
- Zhou, T., & Ji, Y. (2024). On Bayesian sequential clinical trial designs. *The New England Journal of Statistics in Data Science*, 2(1), 136–151. <https://doi.org/10.51387/23-NEJSDS24>

Appendix A: Contents of the Electronic Appendix

The electronic appendix accompanying this thesis contains the source code, the complete experimental results, the original simulation figures, and a README document describing the supplied materials.

The electronic appendix consists of the following items:

- `Appendix_Readme_David_Bartos.pdf`

A guide to the electronic appendix. It contains installation and execution instructions, explains the abbreviated identifiers used in the source code, and describes the organization of the results workbook and the original simulation figures.

- `evolutionary_belief_revision_simulation.py`

The Python source code implementing the belief-update rules, fitness evaluation, evolutionary selection procedure, and diagnostic outputs used in the computational experiments.

- `analysis.xlsx`

The workbook containing the results of all 70 simulation runs, together with the calculations and summary tables used in the analysis presented in Chapter 4.

- `Original simulation figures archive`

A directory containing the original population-composition plots generated during the 70 simulation runs. The directory is divided into seven subdirectories corresponding to the seven parameter configurations described in the methodology.

Each subdirectory contains ten PNG files named `0.0.png` through `0.9.png`. The filename indicates the true coin-bias value used in the corresponding simulation.