
Agent Uncertainty Model and Quantum Mechanics Representation: Non-locality Modeling

Germano Resconi¹ and Boris Kovalerchuk²

¹Dept. of Mathematics and Physics, Catholic University, Brescia, Italy, I-25121
resconi@numerica.it

²Dept. of Computer Science, Central Washington University,
Ellensburg, WA 98926-7520, USA
borisk@cwu.edu

Abstract. This work presents the Agent-based Uncertainty Theory (AUT) and its connection with quantum mechanics where agents are interpreted in terms of the particles. This connection serves a dual goal to justify AUT operations as physically meaningful and to provide a new explanatory mechanism for contradictory issues in quantum mechanics. The AUT is described in agent terms and then is interpreted in quantum mechanics terms. The AUT uses complex aggregations of crisp (non-fuzzy) conflicting judgments of agents. It gives a uniform representation and an operational empirical interpretation for several uncertainty theories such as rough set theory, fuzzy sets theory, intuitionistic fuzzy sets, evidence theory, and probability theory. To build such uniformity the AUT exploits the fact that agents as independent entities can give conflicting evaluations of the same attribute. The AUT many-valued logic is a derived from classical logic with local and global evaluations of proposition p . The evaluation by an individual agent is called a local evaluation and evaluation by a set of agents is called a global evaluation of p . The local operations AND, OR, NOT are classical logic operations, but global operations differ from them. Here a set of agents generates the vectors of logic values. In the AUT local evaluations of proposition p by different agents can be in conflict in contrast with the classical logic. In quantum mechanics, non-locality is related to the superposition of different positions of the particles and to representation of two particles in different positions as a single non-local particle with correlation or entanglement. The logic operations between vectors are the base of AUT many-valued logic combined with the superposition of agents' evaluations.

10.1 Introduction

Complex multi-agent systems (MAS) have many different aspects from the internal structure of the agent to the environment and communications between agents. This paper concentrates on logical aspects of reasoning by individual agents and a society of agents under uncertainty. The common approach in multi-agent modeling in software engineering assumes that agents should be rational to fulfill their mission [3, 14]. Any agent that maximizes chances of success without any logical conflict is considered as a *rational* one. The rationality of a software agent rests in the maximization of its chances of success based on agent's knowledge of its

environment and interaction with the environment and other agents. In fact, an agent's success *criteria*, *reasoning* and *acting* abilities as well as knowledge of the *environment*, can be quite conflicting and uncertain from the logic viewpoint. In such conflicting situation, the agent can be in a logical state that can be called *irrational*. Accordingly, agent's behavior is not rational.

In software engineering, the concept of self-conflict is applicable to the program verification process. If a software agent (program S) produces different results when run several times on the same input data D then S is in a self-conflict state. The reason can be that the non-verified program does not initialize working memory correctly and then reads from memory data C_i that are left from the previous computations.

Another area that is deeply interested in partially rational agents called *boundedly rational* agents is motivated by fundamental goals in *psychology and economics* [5,10]. The works in this area explore the psychology of intuitive beliefs and choices by examining their bounded rationality as Kahneman outlined in his Nobel Prize Lecture in 2002. In economics, bounded rationality means the use of (1) *multiple utility functions* instead of one global scalar utility function, (2) *limited types* of utility functions due to cost of information collection, and (3) *heuristics*.

Our concept of agent in a conflicting logical state (irrational state) is motivated by fundamental goals in the *artificial intelligence* area of logic and reasoning under uncertainty [6,7] with agents [4].

Quantum mechanics is another area where the concept of agents with limited rationality can be useful as we show in this paper. We envision a formal logicization of reasoning of irrational agents that do not follow rational reasoning in the frames of the classical logic and the probability theory and even fuzzy logic. This area is quite different from two previous areas as well as the areas that study expert reasoning simply because irrational agents hardly fit a definition of experts. We define a specific concept of an *ME-rational* (the agent that is rational relative to mutual exclusion) with a clear criterion how to distinguish such agents. Similarly, we define a concept of an *ME-irrational* agent as a *self-conflicting agent* as well as *contradictory agents* that contradict each other but without self-conflict in judgment. The concepts of conflict and self-conflict in judgments are critical for our study of irrational agents, because uncertainty often means that several contradictory statements are made. The concept of self-conflict itself is studied in paraconsistent logics [19] that reason with contradictory statements. Our approach differs from paraconsistent logic in the introduction of agent as fundamental entity for the evaluation.

We build our approach on the results of direct measurements or answers of agents not on the aggregated utility functions to be able to model a structure of contradictions explicitly. Typically current axiomatic theories of uncertainty lack such capability. They assume that initial uncertainty evaluations are already given outside of these theories.

Conducting experiments with devices or asking agents are two common ways to assign initial (basic) uncertainties. Agents are asked to evaluate a proposition p using a classical truth-value logic function $v(p)$ with values in $\{\text{False}, \text{True}\}$. The frequencies of answers provided by a set of agents G are computed. Alternatively each agent is asked to evaluate p using a generalized truth-value logic function $v(p)$ with values in $[0,1]$, where 0 and 1 are classical False and True values and values between 0 and 1 represent intermediate degrees of truth. Both frequency and the

subjective evaluations can be used to get such $v(p)$ as it is done in construction of the fuzzy logic membership function (MF)[7]. In the context of conflicting evaluations of preferences by agents, we discuss the empirical base that includes a mechanism for identifying a type of logic and reasoning computations in the agent logic.

This chapter introduces a hierarchy of logics of uncertainty that is core of the proposed *Agent-based Uncertainty Theory* (AUT). It is a further development of our previous works [8-13] that contain extensive references to related work. The fundamental analysis and review of relevant issues can be found in [1-4, 33-34].

The new contributions include constructing and discovering:

- the agent uncertainty theory (AUT) to model the uncertainty in the society of agents;
- connection among different types of uncertainties by using agents;
- connection between classic logic and fuzzy logic using agents that provides a clear distinction of their domains;
- a new model of many-valued logic and uncertainty by using concepts of conflicts among agents and irrationality of agents;
- a way to model quantum mechanics by using agents and many-valued logic.

The rest of the chapter is structured as follows. Section 10.2 presents the fundamental concepts of conflicts among agents of different orders with the physical interpretation in quantum mechanics. The first order of conflict is defined in section 10.3 Sections 10.4 and 10.5 contain definitions of the conflict of the second order or self-conflict with introduction of agents' correlation in quantum mechanics. In these sections, we define a new negation operator that allows the contradiction to be true and the tautology to be false. Sections 10.6 and 10.7 demonstrate how the self-conflict concept can be used to give a new representation for the evidence theory and the rough sets. This representation helps a user to understand better advantages and disadvantages of these theories and use them more efficiently. Section 10.9 summarises the uniformity of the AUT presentation of different uncertainty theories and section 10.10 describes application areas. Section 10.11 concludes the chapter.

10.2 Concepts and Definitions

The probability calculus does not incorporate explicitly the concepts of irrationality or agent's state of logic conflict. It misses structural information at the level of individual objects, but preserves global information at the level of a set of objects. Given a dice the probability theory studies frequencies of the different faces $E=\{e\}$ as independent (elementary) events. This set of elementary events E has *no structure*. It is only required that elements of E are *mutually exclusive* and *complete*, that is no other alternative is possible. The order of its elements is irrelevant to probabilities of each element of E . No irrationality or conflict is allowed in this definition relative to mutual exclusion. The classical probability calculus does not provide a mechanism for modelling uncertainty when agents communicate (collaborates or conflict). Recent work by Halpern [6] is an important attempt to fill this gap.

Now we will provide more formal definition of AUT concepts. It is done first for individual agents then for sets of agents. Consider a set of agents $G=\{g_1, g_2, \dots, g_n\}$.

Each agent g_k assigns binary true/false value $v \in \{\text{True}, \text{false}\}$ to a proposition p . To show that v was assigned by the agent g_k we will use notation $g_k(p) = v_k$.

Definition. An agent g is called a *reasoning agent* if g assigns a truth-value $v(p)$ to any proposition p from a set of propositions S and any logical formula based on S .

Definition. A set of reasoning agents G is called *totally consistent* for proposition p if any agent g from $G=\{g\}$, always provides the same truth value for p .

In other words, all agents in G are in concord or logical coherence for the same proposition. Thus, changing the agent does not change logic value $v(p)$ for a totally consistent set of agents. Here $v(p)$ is *global* for G and has *no local variability* (independent on the individual agent's evaluation). The classical logic is applicable for such set of consistent (rational) agents.

Definition. A set of reasoning agents G is called *inconsistent* for the proposition p if there are two subset of agents G_1, G_2 such that agents from them provides different truth values for p .

Definition. Let S be a set propositions, $S=\{p_1, p_2, \dots, p_n\}$ then set $\neg S = \{\neg p_1, \neg p_2, \dots, \neg p_n\}$ is called a *complimentary set* of S .

Definition. A set of reasoning agents G is called *S-only-consistent* if agents $\{g\}$ are consistent only for propositions in $S=\{p_1, p_2, \dots, p_n\}$ and are inconsistent in the complimentary set $\neg S$.

The evaluations of p is a vector-function $\mathbf{v}(p)=(v_1(p), v_2(p), \dots, v_n(p))$ for a set of agents G that we will represent as follows:

$$\mathbf{v}(p) = \begin{pmatrix} g_1 & g_2 & \dots & g_{n-1} & g_n \\ v_1 & v_2 & \dots & v_{n-1} & v_n \end{pmatrix} \quad (10.1)$$

An example of the logic evaluation by five agent is shown below for $p="A>B"$

$$f(p) = \begin{pmatrix} g_1 & g_2 & g_3 & g_4 & g_5 \\ true & false & true & false & true \end{pmatrix}$$

Here $A > B$ is true for agents g_1, g_3 , and g_5 and it is false for the agents g_2 , and g_4 .

Kolmogorov's axioms of the probability theory are based on a totally consistent set of agents for a set of statements on mutual exclusion of elementary events. It follows from the following definitions for a set of events $E=\{e_1, e_2, \dots, e_n\}$.

Definition. Set $E=\{e_1, e_2, \dots, e_n\}$ is called a set of *elementary events* (or mutually exclusive events) for predicate E if

$$\forall e_i, e_j \in E \quad E(e_i) \vee E(e_j) = \text{True} \quad \text{and} \quad E(e_i) \wedge E(e_j) = \text{False}.$$

In other words, event e_i is an *elementary event* if for any j , $j \neq i$ events e_i and e_j cannot happen simultaneously, that is probability $P(e_i \wedge e_j) = 0$ and $P(e_i \vee e_j) = P(e_i) + P(e_j)$. Property $P(e_i \wedge e_j) = 0$ is the *mutual exclusion axiom* (ME- axiom).

Let $S = \{p_1, p_2, \dots, p_n\}$ be a set of statements, where $p_i = p(e_i) = \text{True}$ if and only if event e_i is an elementary event. In probability theory, $p(e_i)$ is not associated with any specific agent. It is assumed to be a global property (applicable to all agents). In other words, statements $p(e_i)$ are totally consistent for all agents.

Definition. A set of agent $S = \{g_1, g_2, \dots, g_n\}$ is called a *ME-rational* set of agents if

$$\forall e_i, e_j \in A, \forall g_i \in S, p(e_i) \vee p(e_j) = \text{True} \text{ and } p(e_i) \wedge p(e_j) = \text{False}.$$

In other words, these agents are *totally consistent* or *rational on Mutual Exclusion*. Previously we assumed that each agent assigns value $v(p)$ and we did not model this process explicitly. Now we introduce a set of criteria $C = \{C_1, C_2, \dots, C_m\}$ by which an agent can decide if a proposition p is true or false, i.e., now $v(p) = v(p, C_i)$, which means that p is evaluated by using the criterion C_i .

Definition. Given a set of criteria C , agent g is in a *self-conflicting state* if

$$\exists C_i, C_j (C_i, C_j \in C) \ \& \ v(p, C_i) \neq v(p, C_j)$$

In other words, an agent is in a self-conflicting state if two criteria exist such that p is true for one of them and false for another one.

With the explicit set of criteria, the logic evaluation function $v(p)$ is not vector-function any more, but it is expanded to be a matrix function as shown below:

$$v(p) = \begin{bmatrix} & g_1 & g_2 & \dots & g_n \\ C_1 & v_{1,1} & v_{1,2} & \dots & v_{1,n} \\ C_2 & v_{2,1} & v_{2,2} & \dots & v_{2,n} \\ \dots & \dots & \dots & \dots & \dots \\ C_m & v_{m,1} & v_{m,2} & \dots & v_{m,n} \end{bmatrix} \quad (10.2)$$

For example, four agents using four criteria can produce $v(p)$ as follows:

$$v(p) = \begin{bmatrix} & g_1 & g_2 & g_3 & g_4 \\ C_1 & \text{true} & \text{false} & \text{false} & \text{true} \\ C_2 & \text{false} & \text{false} & \text{true} & \text{true} \\ C_3 & \text{true} & \text{false} & \text{true} & \text{true} \\ C_4 & \text{false} & \text{false} & \text{true} & \text{true} \end{bmatrix}$$

Here agent g_1 is in a self-conflicting state for the criteria C_1 and C_2 , p is true for C_1 and false for C_2 . Similarly, g_3 is in self-conflict for the same criteria. It is also in conflict with g_1 . Agents g_2 and g_4 are not in self-conflict, but in conflict with each other.

If p is a preference statement, $p = (A > B)$ and, $\neg p$ is an opposite preference, $p = (A > B)$ then a self-conflicting agent can provide two or more contradictory preferences, $(p, \neg p)$ for proposition p by stating that p is true and false at the same time.

This can be an indication that g has two or more different competing criteria, $C_1, C_2, \dots, C_m$ to compute the logic value for proposition p . These criteria may or may not be known in advance. The agent may say that car A is better than car B , $p = (A > B) = \text{True}$, and also that B is better than A , $\neg p = (B > A) = \text{True}$. This agent is clearly in self-conflict. It can be a result of using implicitly two different criteria C_1 and C_2 , where C_1 minimizes price and C_2 maximizes quality.

Note that introduction of C explains self-conflict, but does not remove it. The agent still needs to resolve the ultimate preference contradiction having a goal to buy only one and better car. It also does not resolve conflict among different agents if they need to buy a car jointly.

If agent g can modify criteria in C making them consistent for p then g resolves self-conflict. The agent can be in a logic conflict state because of inability to understand the complex *context* and to evaluate *criteria*. For example, in the stock market environment, some traders quite often do not understand a complex context and rational criteria of trading. These traders can be in logic conflict exhibiting chaotic, random, and impulsive behavior. They can sell and buy stocks, exhibiting logic conflicting states “sell” = p and “buy” = $\neg p$ in the same market situation that appears as irrational behavior, which means that the statement $p \wedge \neg p$ can be true.

A logical structure of self-conflicting states and agents is much more complex than it is without self-conflict. In the case of m binary criteria $C_1, C_2, \dots, C_m$ that used to evaluate the logic value for the same attribute, there are 2^m states and only two of them (all true or all false values) do not exhibit conflict between criteria.

10.3 Framework of First Order of Conflict Logic State

10.3.1 Definitions

Definition. A set of agents G is in a first order of conflicting logic state (*first order of conflict*, for short) if

$$\exists g_i, g_j (g_i, g_j \in G) \& v(p, g_i) \neq v(p, g_j).$$

In other words, there are agents g_i and g_j in G for which in

$$v(p) = \begin{pmatrix} g_1 & g_2 & \dots & g_{n-1} & g_n \\ v_1 & v_2 & \dots & v_{n-1} & v_n \end{pmatrix}$$

exist different values v_i and v_j .

Consider n agents g_i that buy cars. There is preference relation or proposition $p = (A > B)$ between cars to be assigned by each of these agents (potential buyers) based on some preference criterion C , $p = \text{True}$ if $A > B$ else $p = \text{False}$. If criterion C is explicitly stated then we can write $A >_C B$. Each agent g_i answers a questionnaire with two options offered for C :

- (1) “ $A > B$ is true” and (2) “ $A > B$ is false”.

Say 70% of agents marked the proposition $p = "A > B \text{ is true}"$, giving $m(A > B) = 0.7$, which measures the level of conflict among n agents. The situation for which a group of non self-conflicting agents assumes that $A > B$ is true and a complementary group of agents assumes that the same $A > B$ is false, is a situation of conflict/contradiction among agents. We denote this situation as a *first order of conflict* for logic statements.

Here each individual agent is completely *rational* and has *no self-conflict*. The conflict exists only *between different agents* when they evaluate the same proposition with only one criterion C . We will write $g(p)$ or $g(A > B)$ to identify that preference statement $p = (A > B)$ is assigned by the agent g . We also will write $g(p, C_k)$ in the case when it is needed to identify a criterion C_k used to compute $g(p)$.

Let $G(A > B)$ be a subset of agents in G such that $A > B = \text{True}$ and $G(A < B)$ be a subset of agents G such that $(A < B) = \text{True}$:

$$G(A > B) = \{g \in G \mid A > B \text{ is true}\}, \quad G(A < B) = \{g \in G \mid A < B \text{ is true}\}.$$

A set of agents G is in the *first order of conflict* if

$$G(A > B) \cap G(A < B) = \emptyset, \quad G(A > B) \neq \emptyset, \quad G(A > B) \cup G(A < B) = G.$$

The following definition presents this idea in general terms.

Definition. A set of agents G is in the *First Order Conflict (FOC)* for proposition p if

$$G(p) \cap G(\neg p) = \emptyset, \quad \text{and } G(p) \neq \emptyset, \quad G(p) \cup G(\neg p) = G.$$

Example. Let $G = \{g_1, g_2, g_3, g_4\}$ and for two agents $G(p) = \{g_1, g_4\}$ from G proposition p is true and for two other agents $G(\neg p) = \{g_2, g_3\}$ proposition $\neg p$ is true, and p is false where $p = (A > B)$. It can be written as $g_1(p) = \text{True}$, $g_2(p) = \text{False}$, $g_3(p) = \text{false}$, $g_4(p) = \text{True}$. The set of agents G is in a conflict logic state. We record the vector of logic values of p provided by all four agents as

$$\mathbf{v}(p) = \begin{bmatrix} g_1 & g_2 & g_3 & g_4 \\ v(p) = \text{true} & v(p) = \text{false} & v(p) = \text{false} & v(p) = \text{true} \end{bmatrix}$$

Fig. 10.1 shows a set of 20 agents in the logic conflicting state, where 7 white agents are in the state True and 13 black agents are in the state False for the same proposition $p = "A > B"$.

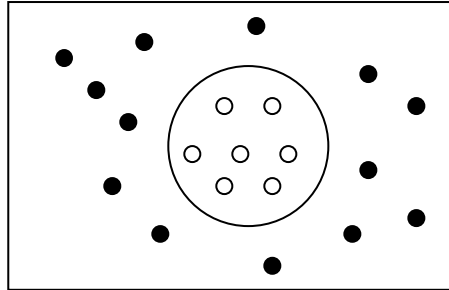


Fig. 10.1. A set of 20 agents in the first order of logic conflict

Below we show that at the first order of conflicts, AND and OR operations should not be the traditional classical logic operations, but should be *vector operations* in the space of the agents' evaluations (*agents space*). The vector operations reflect a structure of logic conflict among coherent individual agent evaluations.

10.3.2 Fusion Process

If a single decision must be made at the first order of conflict, then we must introduce a *fusion process* of the logic values of the proposition p given by all agents. A basic way to do this is to compute the weighted frequency of logic value given by all agents:

$$\mu(p) = w_1 v_1(p) + \dots + w_n v_n(p) = \begin{bmatrix} v_1(p) \\ v_2(p) \\ \dots \\ v_n(p) \end{bmatrix}^T \begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \end{bmatrix} \quad (10.3)$$

where $\begin{bmatrix} v_1(p) \\ v_2(p) \\ \dots \\ v_n(p) \end{bmatrix}$ and $\begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \end{bmatrix}$ are two vectors in the space of the agents' evaluations.

The first vector contains all logic states (True/False) for all agents, the second vector (with property $\sum_{k=1}^n w_k = 1$) contains non-negative weights (utilities) that are given to each agent in the fusion process. In a simple frequency case, each weight is equal to $1/n$.

At first glance, $\mu(p)$ is the same as used in the probability and utility theories. However, classical axioms of the probability theory have no references to agents producing initial uncertainty values and do not violate the mutual exclusion. As a result, formulas for assigning initial uncertainty values such as $\mu(A > B)$ are not a part of the theory. Value $\mu(p)$ can differ significantly from the classical relative frequency for some irrational agents.

A probability value can be agent's judgment or a result of physical experiments. In [2] agents are connected to logic and probability in the following way. In a given state s , the formula $P_j(\varphi)$ denotes the probability of the logic proposition φ according to the probability distribution for agent g_j in the state s . In this model, an individual agent gives the probability. All agents are autonomous and no conflict among agents and self-conflict is modeled explicitly. This formalization does not provide tools to fuse the contradictory knowledge of different agents that we described.

How we can justify formula (3) beyond the reference to the utility and probability theories? It can be done in two general ways by using concepts of descriptive and prescriptive theories to clarify this issue. To be descriptive (3) should have weight $\{w_i\}$ that model a specific task at hand and properly instantiated using available data.

To be prescriptive (3) must be inferred in a more general way, say, from axioms that are appropriate for a particular task. Below we define *vector logic operations* for the first order of conflict logic states $\mathbf{v}(p)$.

Consider $\mu_{\max} = \max (\mu(p))$. This maximum is reached for p such that $v_i(p)=1$ for all agents g_i . If in addition to this $\sum_{k=1}^n w_k=1$ and $w_k \geq 0$ then $\mu_{\max}(p)=1$.

Definition. The *Agent Set Contradiction (ASC) index* $\lambda(p)$ for proposition p is defined as follows

$$\lambda(p) = \begin{cases} 1 - \mu(p) / \mu_{\max}(p), & \text{if } \mu(p) \geq 0.5 \mu_{\max}(p) \\ \mu(p) / \mu_{\max}(p), & \text{if } \mu(p) < 0.5 \mu_{\max}(p) \end{cases}$$

For instance, if $\mu(p) = 0.7$ then $\lambda(p) = 1 - \mu_{\max}(p) = 0.3$ if $\mu_{\max}(p) = 1$ and if $\mu(p) = 0.3$ then $\lambda(p) = \mu(p) = 0.3$. Thus in both cases, $\lambda(p)$ is the same showing the equal difference from the certain (True, False) values.

Definition. Agent g is *irrational for proposition* p if $Ir(g,p) = 1$, where $Ir(g,p) = 1 \Leftrightarrow g(p) \wedge g(\neg p) = \text{True}$ or $g(p) \vee g(\neg p) = \text{False}$.

Definition. The *Individual Agent Irrationality (IAI) index* $\gamma(g,P)$ for a set of h propositions P is defined as follows:

$$\gamma(g, P) = \frac{1}{h} \sum_{i=1}^h Ir(g, p_i)$$

where $Ir(g,p) = 1 \Leftrightarrow g(p) \wedge g(\neg p) = \text{True}$ or $g(p) \vee g(\neg p) = \text{False}$.

For instance, if $\gamma(g,P) = 0.5$ then for 50% of propositions agent g exhibits irrational behavior. In this case it is better to avoid using the standard probabilistic technique, but if $\gamma(g,P) = 0.0001$ then it can be very reasonable to use the probabilistic technique for propositions P for this individual agent.

Definition. Agent g is *irrational for proposition* p and *criterion* C_k if $Ir(g,p,C_k) = 1$, where $Ir(g,p,C_k) = 1 \Leftrightarrow g(p,C_k) \wedge g(\neg p,C_k) = \text{True}$ or $g(p,C_k) \vee g(\neg p,C_k) = \text{False}$.

Definition. The *Individual Agent Irrationality (IAI) index* $\gamma(g,p,C)$ for proposition p and m criteria C is defined as follows:

$$\gamma(g, p, C) = \frac{1}{m} \sum_{k=1}^m Ir(g, p, C_k)$$

where $Ir(g,p,C_k) = 1 \Leftrightarrow g(p,C_k) \wedge g(\neg p,C_k) = \text{True}$ or $g(p,C_k) \vee g(\neg p,C_k) = \text{False}$.

For instance, if $\gamma(g,P) = 0.5$ then a half of the agents are irrational at some extend, that is for 50% of agents at least one proposition exists for which g exhibits irrational behavior. If $\gamma(p,C) = 0.5$ then

For instance, if $\gamma(g,p,C)=0.5$ then for 50% of criteria agent g exhibits irrational behavior for proposition p . This means that explicit identification of the criterion C_k does not resolve the contradiction for 50% of criteria. This is the indication that 50% of criteria do not capture the source of the irrationality specifically enough and further specification of criteria can be needed before a technique that is applicable to rational agents will be used. If $\gamma(g,p,C)=0.001$ then agent g has logic conflict with p only for 1% of criteria. Often this 1% can be ignored and the probabilistic technique can be used.

Definition

$$\begin{aligned} v(p \wedge q) &= v_1(p) \wedge v_1(q), \dots, v_n(p) \wedge v_n(q), \\ v(p \vee q) &= v_1(p) \vee v_1(q), \dots, v_n(p) \vee v_n(q) \\ v(\neg p) &= \neg v_1(p), \dots, \neg v_n(p), \end{aligned}$$

where the symbols $\wedge, \vee, \neg$ in the right side of the equations are the classical AND, OR, and NOT operations.

These operations can be written also with explicit indication of agents involved in the first row:

$$\begin{aligned} v(p \wedge q) &= \begin{pmatrix} g_1 & g_2 & \dots & g_{n-1} & g_n \\ v_1(p) \wedge v_1(q) & v_2(p) \wedge v_2(q) & \dots & v_{n-1}(p) \wedge v_{n-1}(q) & v_n(p) \wedge v_n(q) \end{pmatrix} \\ v(p \vee q) &= \begin{pmatrix} g_1 & g_2 & \dots & g_{n-1} & g_n \\ v_1(p) \vee v_1(q) & v_2(p) \vee v_2(q) & \dots & v_{n-1}(p) \vee v_{n-1}(q) & v_n(p) \vee v_n(q) \end{pmatrix} \\ v(\neg p) &= \begin{pmatrix} g_1 & g_2 & \dots & g_{n-1} & g_n \\ v_1(\neg p) & v_2(\neg p) & \dots & v_{n-1}(\neg p) & v_n(\neg p) \end{pmatrix} \end{aligned}$$

Below we present the important properties of sets of conflicting agents at the first order of conflicts. Let $|G(x)|$ be the numbers of agents for which x is true.

Statement 1. If G is a set of agents at the first order of conflicts and

$$\begin{aligned} |G(q)| &\leq |G(p)| \quad \text{then} \\ |G(p \wedge q)| &= \min(|G(p)|, |G(q)|) - |G(\neg p \wedge q)|, \end{aligned} \tag{10.4}$$

$$|G(p \vee q)| = \max(|G(p)|, |G(q)|) + |G(\neg p \wedge q)|. \tag{10.5}$$

If G is a set of agents at the first order of conflicts and

$$\begin{aligned} |G(p)| &\leq |G(q)| \quad \text{then} \\ |G(p \wedge q)| &= \min(|G(p)|, |G(q)|) - |G(\neg q \wedge p)|, \\ |G(p \vee q)| &= \max(|G(p)|, |G(q)|) + |G(\neg q \wedge p)|. \end{aligned}$$

and also $|G(p \vee q)| = |G(p) \cup G(q)|$, $|G(p \wedge q)| = |G(p) \cap G(q)|$

Below we illustrate correctness of Statement 1 with two examples for three agents. In the example (a) three agents g_1, g_2 , and g_3 provided answers $v_1(p)=1, v_2(p)=v_3(p)=0$ and $v_1(q)=v_2(q)=1$ and $v_3(q)=0$, respectively.

In the example (b) the answers are different: $v_1(p)=0, v_2(p)=v_3(p)=1$, and $v_1(q)=v_2(q)=0, v_3(q)=1$.

In example (a) only one agent (g_1) answered that p is true, thus $|G(p)|=1$. For q two agents (g_1, g_2) answered that q is true, thus $|G(q)|=2$ and $|G(p)| < |G(q)|$. Therefore, formula (4) is not applicable to this example. Its assumption $|G(q)| < |G(p)|$ is false.

In example (b) two agents (g_2, g_3) answered that p is true, thus $|G(p)|=2$. For q one agent (g_3) answered that q is true, thus $|G(q)|=1$. Hence, $|G(q)| < |G(p)|$ and (4) can be applied for this example. Its assumption is true. Here $\min(|G(p)|, |G(q)|) = 1$ and $|G(\neg p \wedge q)| = 0$ because $g_1(\neg p \wedge q) = 0, g_2(\neg p \wedge q) = 0, g_3(\neg p \wedge q) = 0$.

Thus, $\min(|G(p)|, |G(q)|) - |G(\neg p \wedge q)| = 1$. We have $|G(p \wedge q)| = 1$ on the left side of (4) because $g_1(p \wedge q)=0, g_2(p \wedge q)=0, g_3(p \wedge q)=1$. Thus, statement (1) is correct for example (b).

Corollary 10.1. If G is a set of agents at the first order of conflicts such that $G(q) \subset G(p)$ or $G(p) \subset G(q)$ then

$$\begin{aligned} G(\neg p \wedge q) &= \emptyset \text{ or } G(\neg q \wedge p) = \emptyset \\ |G(p \wedge q)| &= \min(|G(p)|, |G(q)|) \\ |G(p \vee q)| &= \max(|G(p)|, |G(q)|) \end{aligned}$$

This follows from the statement 1. The corollary presents a well-known condition when the use of min, max operations has the clear justification.

Let $G^c(p)$ is a *complement* of $G(p)$ in G : $G^c(p) = G \setminus G(p)$, $G = G(p) \cup G^c(p)$.

Statement 2. $G = G(p) \cup G^c(p) = G(p) \cup G(\neg p)$.

Corollary 10.2. $G(\neg p) = G^c(p)$. It follows directly from Statement 2.

Statement 3. If G is a set of agents at the first order of conflicts then

$$\begin{aligned} G(p \vee \neg p) &= G(p) \cup G(\neg p) = G(p) \cup G^c(p) = G \\ G(p \wedge \neg p) &= G(p) \cap G(\neg p) = G(p) \cap G^c(p) = \emptyset \end{aligned}$$

It follows from the definition of the first order of conflict and statement 2. In other words, $G(p \wedge \neg p) = \emptyset$ corresponds to the contradiction $p \wedge \neg p$, that is always false and $G(p \vee \neg p) = G$ corresponds to the tautology $p \vee \neg p$, that is always true in the first order conflict.

Let $G_1 \oplus G_2$ be a *symmetric difference* of sets of agents G_1 and G_2 ,

$$G_1 \oplus G_2 = (G_1 \cap G_2^c) \cup (G_1^c \cap G_2)$$

and let $p \oplus q$ be the *exclusive or* of propositions p and q ,

$$p \oplus q = (p \wedge \neg q) \vee (\neg p \wedge q).$$

Consider, a set of agents $G(p \oplus q)$. It consists of agents for which values of p and q . Below we use the number of agents in set $G(p \oplus q)$ to define a differ from each other, that is

$$G(p \oplus q) = G((p \wedge \neg q) \vee (\neg p \wedge q)).$$

measure of difference between statements p and q and a measure of difference between sets of agents $G(p)$ and $G(q)$.

Definition. *Measure of difference* $D(p, q)$ between statements p and q and a measure of difference $D(G(p), G(q))$ between sets of agents $G(p)$ and $G(q)$ are defined as follows:

$$D(p, q) = D(G(p), G(q)) = |G(p) \oplus G(q)|$$

Statement 4. $D(p, q) = D(G(p), G(q))$ is a distance, i.e., it satisfies distance axioms

$$D(p, q) \geq 0$$

$$D(p, q) = D(q, p)$$

$$D(p, q) + D(q, h) \geq D(p, h).$$

This follows from the properties of the symmetric difference $\oplus$ [Flament 1963].

Fig. 10.2 illustrates a set of agents $G(p)$ for which p is true and a set of agents $G(q)$ for which q is true. In Fig. 10.2(a) the number of agents for which truth values of p and q are different, $(\neg p \wedge q) \vee (p \wedge \neg q)$, is equal to 2. These agents are represented by white squares. Therefore the distance between $G(p)$ and $G(q)$ is 2. Fig. 10.2(b) shows other $G(p)$ and $G(q)$ sets with the number of the agents for which $\neg p \wedge q \vee (p \wedge \neg q)$ is true equal to 6 (agents shown as white squares and squares with the grid). Thus, the distance between the two sets is 6.

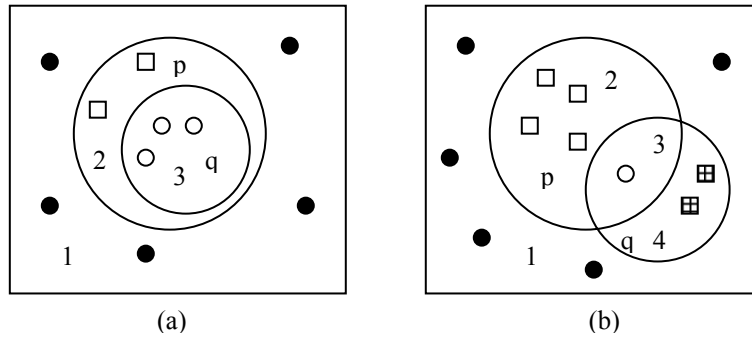


Fig. 10.2. A set of total 10 agents with two different splits to $G(p)$ and $G(q)$ subsets (a) and (b). These splits produce different distances between $G(p)$ and $G(q)$. The distance in the case (a) is equal to 2, the distance in the case (b) is equal to 6. Set 1 (black circles) consists of agents for which both p and q are false, Set 2 (white squares) consists of agents for which p is true but q is false. Set 3 (white circles) consists of agents for which p and q are true, and Set 4 (squares with grids) consists of agents for which p is false and q is true.

In Fig. 10.2(a), set 2 consists of 2 agents $|G((p \wedge \neg q))| = 2$ and set 4 is empty, $|G((\neg p \wedge \neg q))| = 0$, thus $D(\text{Set2}, \text{Set4})=2$. This emptiness means that a set of agents with true p includes a set of agents with true q . In Fig. 10.2(b), set 2 consists of 4 agents $|G((p \wedge \neg q))| = 4$ and set 4 consists of 2 agents, $|G((\neg p \wedge \neg q))| = 2$, thus $D(\text{Set2}, \text{Set4})=6$.

If $G(p)$ includes $G(q)$ or $G(q)$ includes $G(p)$ (case (a) in Fig. 10.2) then

$$\begin{aligned} |G(p \wedge q)| &= |G(p) \cap G(q)| = \min(|G(p)|, |G(q)|) = 3 = |G(q)| \\ |G(p \vee q)| &= |G(p) \cup G(q)| = \max(|G(p)|, |G(q)|) = 5 = |G(p)| \end{aligned}$$

In this case nothing is added or subtracted in (4) and (5),

$$\begin{aligned} |G(p \wedge q)| &= \min(|G(p)|, |G(q)|) - |G(\neg p \wedge q)|, \\ |G(p \vee q)| &= \max(|G(p)|, |G(q)|) + |G(\neg p \wedge q)|, \end{aligned}$$

Therefore, $|G(p \wedge q)|$ reaches its maximum value and $|G(p \vee q)|$ reaches its minimum value. In the case (b) in Fig. 10.2 the situation is different with the distance equal to 6, $|G(\neg p \wedge q)| = 2$ and

$$\begin{aligned} |G(p \wedge q)| &= |G(p) \cap G(q)| = \min(|G(p)|, |G(q)|) - |G(\neg p \wedge q)| = 1 \\ |G(p \vee q)| &= |G(p) \cup G(q)| = \max(|G(p)|, |G(q)|) + |G(\neg p \wedge q)| = 7 \end{aligned}$$

If sets of agents do not overlap then the distance between them reaches its maximum on these sets. Similarly, for AND operation the value of $|G(p \wedge q)|$ reaches its minimum and for OR operation the value of $|G(p \vee q)|$ reaches its maximum,

$$\begin{aligned} |G(p \wedge q)| &= |G(p) \cap G(q)| = \min(|G(p)|, |G(q)|) - |G(\neg p \wedge q)| = 0 \\ |G(p \vee q)| &= |G(p) \cup G(q)| = \max(|G(p)|, |G(q)|) + |G(\neg p \wedge q)| = |G(p)| + |G(q)| \end{aligned}$$

Therefore, with the same $|G(p)|$ and $|G(q)|$ we may have different distances and different results for AND and OR operations among the sets. These operations depend on the internal structures of the agents' states. They are not simple truth-functional operations of the classical propositional logic. We remark that in this way t-norm and t-conorm (that are truth-functional) can be redefined to become context dependent.

The negation operator “ $\neg$ ” for agents satisfies the De Morgan rule:

$$\begin{aligned} G[\neg(p \wedge q)] &= G(p \wedge q)^C = [G(p) \cap G(q)]^C = G[\neg p \vee \neg q] = \\ G(\neg p) \cup G(\neg q) &= G^C(p) \cup G^C(q) = \max(G(\neg p), G(\neg q)) + G(\neg p \wedge \neg q) \\ G[\neg(p \vee q)] &= G(p \vee q)^C = G[\neg p \wedge \neg q] = G(\neg p) \cap G(\neg q) = \\ G(p)^C \cap G(q)^C &= \min(G(\neg p), G(\neg q)) - G(\neg p \wedge \neg q), \end{aligned}$$

and it can be used for reasoning with agents' logic evaluations that involves negation.

For complex computations of logic values provided by the agents we can use a *graph of the distances* among the sentences $p_1, p_2, \dots, p_N$. For example, for three sentences p_1, p_2 , and p_3 , we have a graph of the distances shown in Fig. 10.3.

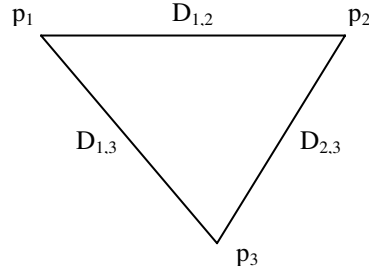


Fig. 10.3. Graph of the distances for three sentences p_1 , p_2 , and p_3

This graph has can be represented by a distance matrix

$$D = \begin{bmatrix} 0 & D_{1,2} & D_{1,3} \\ D_{1,2} & 0 & D_{2,3} \\ D_{1,3} & D_{2,3} & 0 \end{bmatrix}$$

which has a general form of a symmetric matrix,

$$D = \begin{bmatrix} 0 & D_{1,2} & \dots & D_{1,N} \\ D_{1,2} & 0 & \dots & D_{2,N} \\ \dots & \dots & \dots & \dots \\ D_{1,N} & D_{2,N} & \dots & 0 \end{bmatrix}$$

Having the distances D_{ij} between propositions we can use them to compute complex expressions in agents' logic operation.

For instance, using $0 \leq D(p, q) \leq G(p) + G(q)$ and

$$G(p \wedge q) \equiv G(p \vee q) - G((\neg p \wedge q) \vee (p \wedge \neg q)) \equiv G(p \vee q) - D(p, q)$$

we can infer

$$D(p, q) = 0 \Rightarrow G(p \wedge q) = G(p \vee q) \quad (10.6)$$

and

$$D(p, q) = G(p) + G(q) \Rightarrow G(p \vee q) = G(p) + G(q) \text{ and } G(p \wedge q) = 0 \quad (10.7)$$

Similarly, we can infer:

$$G(\neg p \wedge q) \equiv D(p, q) - G(p \wedge \neg q) \quad (10.8)$$

for expression $\neg p \wedge q$,

$$G(\neg p \vee q) \equiv G(\neg p \wedge q)^c = N - (D(p, q) - G(p \wedge \neg q)) \quad (10.9)$$

for the implication rule $\neg p \vee q = p \rightarrow q$,

$$G(q) = G(p_1 \vee p_2) - D(p_1, p_2) \quad (10.10)$$

$$G(F) = G(q \vee p_3) - D(q, p_3) \quad (10.11)$$

for $F = p_1 \wedge p_2 \wedge p_3$ and $q = (p_1 \wedge p_2)$.

In this section, we formalized the concept of first order conflicting agents with the conflict only between different agents while each agent has no self-conflict. It is shown that for these conflicts AND and OR operations should be defined as vector operations in the agent space. Such operations preserve the coherence of individual agent evaluations. Next, we had shown that evaluations for the first order conflicting agents can differ from traditional evaluations that do not consider conflict among agents.

10.3.3 Quantum Mechanics Superposition and First Order Conflicts

The Mutual Exclusive (ME) principle states that an object cannot be in two different positions at the same time. In quantum mechanics, we have *non-local phenomena* for which the same particle can be in many different positions at the same time. This is a clear violation of the ME principle. The non-locality is essential for quantum phenomena. Given proposition $p =$ “The particle is in position x in the space”, agent g_i can say that p is true but agent g_j can say that the same p is false.

Individual states of the particle include its position, momentum and others. The *complete state* of the particle is a *superposition* of the quantum states of the particle w_i at different positions x_i in the space, This superposition can be a global wave function:

$$\psi = w_1 |x_1\rangle + w_2 |x_2\rangle + \dots + w_n |x_n\rangle,$$

where $|x_i\rangle$ denotes the state at the position at x_i [26]. If we limit our consideration by an individual quantum state of the particle and ignore other states then we are in Mutual Exclusion situation of the classical logic (the particle is at x_i or not). In the same way in AUT, when we observe only one agent at the time the situation collapses to the classical logic that may not adequately represent many-valued logic properties. Having multiple states associated with the particle, we use AUT multi-valued logic as a way to model the multiplicity of states. This situation is very complex because measuring position x_i of the particle changes the superposition value ψ . However, this situation can be modeled by the first order conflict because each individual agent has no self-conflict.

10.4 Second Order of Conflicts, Self-conflicts, and Irrationality

10.4.1 Definitions and Examples

At the first order of logic conflict, we have $G(\neg p) = G^C(p)$ and fused $\mu(\neg p)$ satisfies the following statement 5.

Statement 5

$$\mu(\neg p) = w_1 v_1(\neg p) + \dots + w_n v_n(\neg p) =$$

$$\begin{bmatrix} v_1(\neg p) \\ v_2(\neg p) \\ \dots \\ v_n(\neg p) \end{bmatrix}^T \begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \end{bmatrix} = \begin{bmatrix} 1-v_1(p) \\ 1-v_2(p) \\ \dots \\ 1-v_n(p) \end{bmatrix}^T \begin{bmatrix} w_1 \\ w_2 \\ \dots \\ w_n \end{bmatrix} = 1 - \mu(p)$$

This property is consistent with the known Zadeh's property $\mu(\neg p) + \mu(p) = 1$. Atanassov assumes in his model [16] three entities: property p , its negation $\neg p$, and an indeterminacy s (about p or its negation) with the following definition of their relations,

$$\mu(\neg p) + \mu(p) + \mu(s) = 1.$$

Atanassov's model is also a *fuzzy partition* as defined by Ruspini [24]. Atanassov's model differs from Lukasiewicz's three-value logic that has a fixed $\mu(s) = 0.5$, which is interpreted as possibility [15]. The simple observation suggests a potential conflict between Zadeh's and Atanassov's models noted in [15]. To give a meaning to Atanassov's model we introduce the second order of conflict or the *logical self-conflict* (LSC) among agents into this model.

We remark that at the first order of conflict the distance between p and $\neg p$ is

$$D(p, \neg p) = |G(p)| + |G(\neg p)| = n,$$

where n is the total number of agents. Conceptually it follows from the fact that in the first order of conflict every agent uses only one criterion C to obtain the logic state true or false for the proposition p . In contrast, when an agent uses two criteria C_1, C_2 to obtain the logic value for p we can get more complex situations, e.g.:

	C_1	C_2
$State_1$	true	true
$State_2$	true	false
$State_3$	false	true
$State_4$	false	false

The states 1 and 4 are internally coherent states, where the agent is not in a self-conflict. In state 1, both criteria give true value, and in state 4, both criteria give false value for proposition p . The states 2 and 3 are self-conflicting logic states. In state 2, Criterion C_1 gives p true value, but criterion C_2 for the same p gives false value. Therefore, the same proposition is true and false in the same time and this generates agent's self-conflict. For state 3, the situation is opposite, where the agent gives the logic value false for C_1 and true for C_2 . Here for the same proposition p we have two different logic values given by the same agent. It is clear that this is again a self-conflict situation. Thus, for two criteria we have four states of agents. An example is shown in Fig. 10.4.

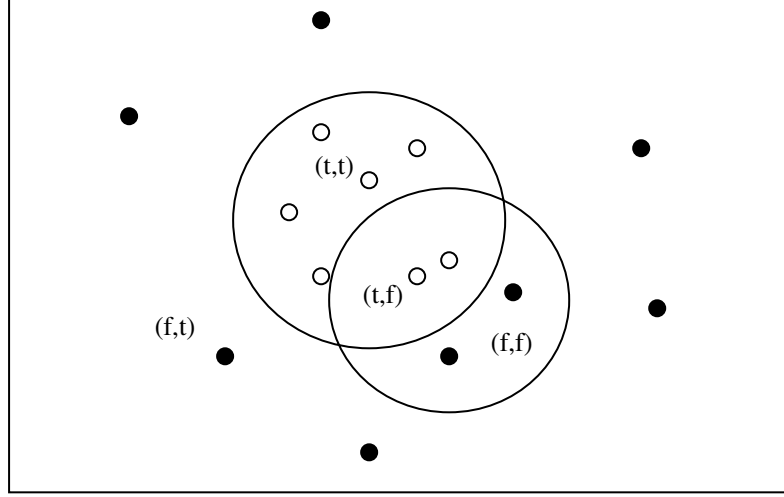


Fig. 10.4. Illustration of two criteria C_1 and C_2 . White circles are agents for which p is true for criterion C_1 and black circles are agents for which p is false and $\neg p$ is true for C_1 .

Let agents $g_1, g_2, \dots, g_n$ be agents in G that use criteria C_1 and C_2 to assign logic values $v_{1,j}$ and $v_{2,j}$, respectively, which are true or false.

Definition. An agent g is a *self-conflicting agent* if $v(g,p,C_1) \neq v(g,p,C_2)$, where $v(g,p,C_1)$ and $v(g,p,C_2)$ are values assigned by agent g to proposition p using criterion C_1 or C_2 , respectively.

Any self-conflicting agent can be denoted as an *irrational agent*. If p is preference relation and criteria C_1 and C_2 are known then the conflict is explicit and it has the explanation. It can be a conflict between cost (C_1) and reliability (C_2). If C_1 and C_2 are not known explicitly then the conflict is *implicit* and it has no explanations by criteria.

Definition. A set of agents G is in *the second order of conflict* if a self-conflicting agent g exists in G .

Formally, the complete evaluation by all agents for the two criteria is a matrix:

$$\mathbf{v}(p) = \begin{bmatrix} g_1 & g_2 & \dots & g_n \\ C_1 & v_{1,1} & v_{1,2} & v_{1,n} \\ C_2 & v_{2,1} & v_{2,2} & v_{2,n} \end{bmatrix}$$

The first row indicates agents, the second row contains the logic value for the first criterion, and the third row contains the logic value for the second criterion.

Below we show that at the second order of conflicts, the classical negation takes place if the agent uses one criterion at the time. Let agent g_k accepts proposition p as true using criterion C_1 . In this case, two evaluations are possible when g_k uses C_2 :

$$\mathbf{v}_1(p) = \begin{bmatrix} g_k \\ true \\ false \end{bmatrix}, \mathbf{v}_2(p) = \begin{bmatrix} g_k \\ true \\ true \end{bmatrix}$$

If agent g_k evaluates p using C_1 as false then we have two other cases:

$$\mathbf{v}_3(p) = \begin{bmatrix} g_k \\ false \\ false \end{bmatrix}, \mathbf{v}_4(p) = \begin{bmatrix} g_k \\ false \\ true \end{bmatrix}$$

If g_k ignores values provided by C_2 when negating results provided by using C_1 then it is classical logic situation to produce negation as shown below for $\mathbf{v}_3(\neg p)$ and $\mathbf{v}_4(\neg p)$:

$$\mathbf{v}_3(\neg p) = \begin{bmatrix} g_k \\ true \\ false \end{bmatrix}, \mathbf{v}_4(\neg p) = \begin{bmatrix} g_k \\ true \\ true \end{bmatrix}$$

Again, in this situation the second criterion C_2 is completely *independent* from the first criterion C_1 and has no influence on the logic evaluation of p using C_1 . In a graphic form, this situation was depicted in Fig. 10.4. If g_k evaluated p as true using the second criterion C_2 then two cases are possible:

$$\mathbf{v}_5(p) = \begin{bmatrix} g_k \\ false \\ true \end{bmatrix}, \mathbf{v}_6(p) = \begin{bmatrix} g_k \\ true \\ true \end{bmatrix}$$

Again, assuming here that agent g_k ignores the first criterion when negating the truth value produced by C_2 we get:

$$\mathbf{v}_5(\neg p) = \begin{bmatrix} g_k \\ false \\ false \end{bmatrix}, \mathbf{v}_6(\neg p) = \begin{bmatrix} g_k \\ true \\ false \end{bmatrix}$$

In a graphic form, it is shown in Fig. 10.5.

Now when p is true for the first criteria and false for the second criteria, we obtain

$$\mathbf{v}(p) = \begin{bmatrix} g_k \\ C_1 & true \\ C_2 & true \end{bmatrix}, \mathbf{v}(p) = \begin{bmatrix} g_k \\ C_1 & true \\ C_2 & false \end{bmatrix}$$

$$\mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & true \\ C_2 & false \end{bmatrix}, \mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & false \\ C_2 & false \end{bmatrix}$$

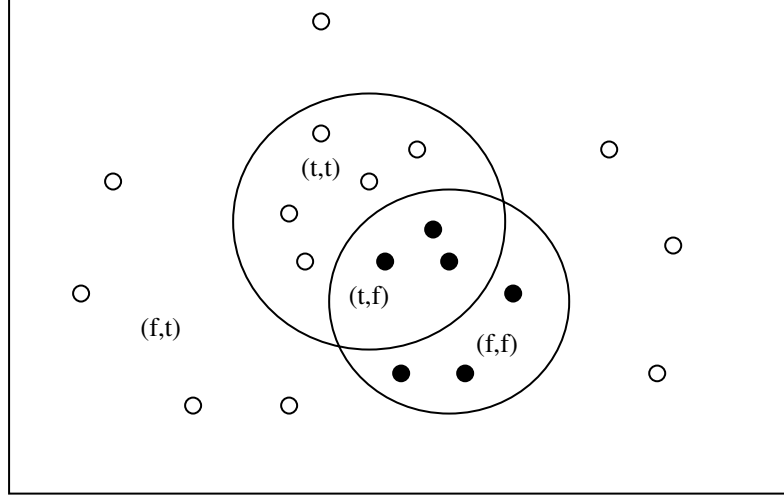


Fig. 10.5. Illustration of independent criterion C_2 . White circles are agents for which p is true and black circles are agents for which p is false and $\neg p$ is true for criterion C_2 .

The agents in self-conflict states belong to the set of agents for which p is true and p is not true. Thus,

$$\mathbf{v}(p) \wedge \mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & \text{true} \\ C_2 & \text{false} \end{bmatrix} \wedge \begin{bmatrix} g_k \\ C_1 & \text{true} \\ C_2 & \text{false} \end{bmatrix} = \begin{bmatrix} g_k \\ C_1 & \text{true} \\ C_2 & \text{false} \end{bmatrix}$$

The contradiction is true for the first criterion, but when p is false for the first criterion and true for the second one, we have

$$\mathbf{v}(p) = \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{true} \end{bmatrix}, \mathbf{v}(p) = \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{false} \end{bmatrix}$$

$$\mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & \text{true} \\ C_2 & \text{true} \end{bmatrix}, \mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{true} \end{bmatrix}$$

and

$$\mathbf{v}(p) \wedge \mathbf{v}(\neg p) = \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{true} \end{bmatrix} \vee \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{true} \end{bmatrix} = \begin{bmatrix} g_k \\ C_1 & \text{false} \\ C_2 & \text{true} \end{bmatrix}$$

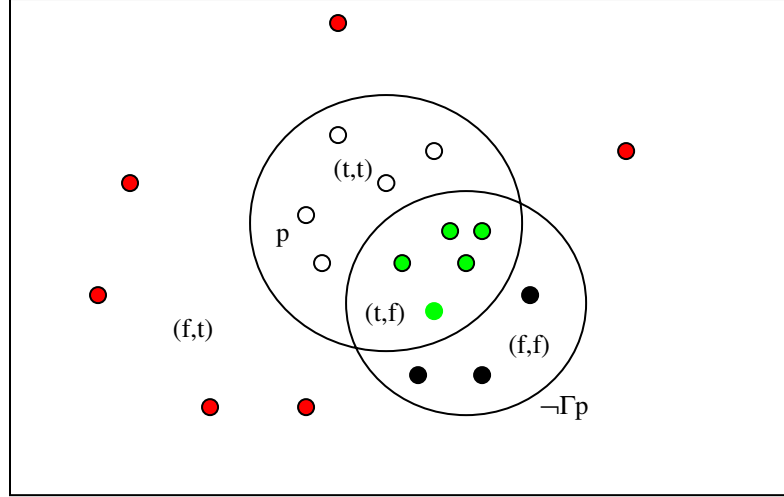


Fig. 10.6. Example with two criteria. The white circles are the agents for which p is true for both C_1 and C_2 criteria, which is denoted as (t,t) . The black circles are the agents for which p is false for both C_1 and C_2 criteria, that is $\neg p$ is true for C_1 and C_2 . The green agents are self-conflicting agents for which the contradiction $p \wedge \neg p$ is true for the first criterion C_1 . For the red self-conflicting agents the tautology is false for the first criteria

The tautology is false for the first criterion. Fig. 10.6 provided another example explained in its caption.

We can remark that the graph in Fig. 10.6 can be obtained by the superposition of the two criteria C_1 , C_2 as we show in Fig. 10.7.

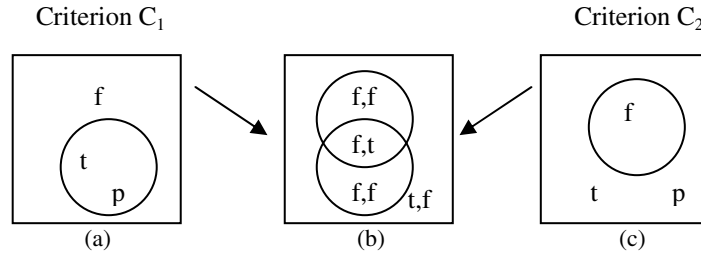


Fig. 10.7. Criteria and logic values in the second order of conflict. Case (a) presents a set where p is true for the first criterion, and case (c) presents a set where p is true for the second criterion. Case (b) shows the superposition of the two criteria.

10.4.2 Correlation in Quantum Mechanics and Second Order Conflict

Consider two interacting particles as agents. These agents are interdependent. This correlation is independent from any possible physical communication by any type of fields. We can say that the correlation is more a logic state for the particles that are

dependent. Now the quantum correlation is a conflicting state because we know from quantum mechanics that there is a correlation but when we try to measure the correlation we cannot check the correlation itself. The measure destroys the correlation. So if the spin of one electron is up and another electron is down the change of the first spin when we have correlation or entanglement generate the change of the other spin instantaneously and this is manifested in a statistic correlation different from zero.

10.5 Agents and Atanassov's - Ruspini's Models

This section shows that Atanassov's intuitionistic fuzzy sets [22] and Ruspini's model [24] can gain an interpretation using agents' approach. There is a following property in [22]:

$$G(p) \cup G(\neg p) + G(1) = \text{Universe}$$

Now consider an example in Fig. 10.8, where $G(p)$ is a set of agents in areas (2) and (3), $G(\neg p)$ is a set of agents in areas (3) and (4), and $G(1)$ is the set of agents in area (1) with a vector of values (f, t) .

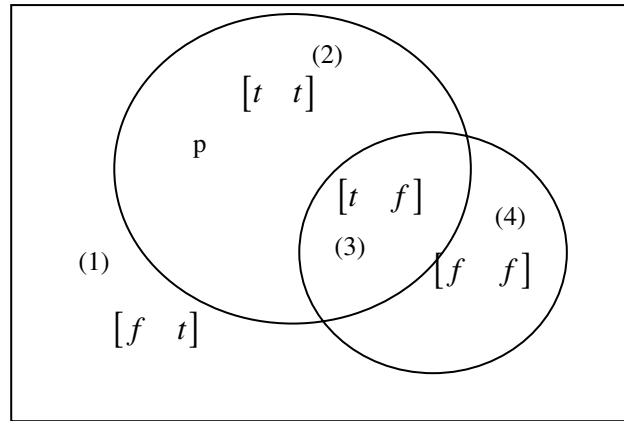


Fig. 10.8. Four classes of agents. The logic value true for p is located in $(2, 3)$, the criteria negation of p is located in $(3, 4)$. Set (3) and set (1) are the logically self-conflicting states.

Let set (2,3) be "Tall" and set (4,3) be "Short", that is different from the word "not tall" which is set (1,4). This consideration is in line with Atanassov's work. Now with the *distance* among set of agents that represent propositions we can define the distance between sets "Tall" and "Short". This gives agents' interpretation of the Atanassov's theory. In set 3, we have the superposition of p and negation $\neg p$, where $\neg p$ is true for C_2 . In sets 3 and 4, p is false for the criterion C_2 and p is true. In sets 2 and 3, p is true for criterion C_1 . In addition, in set 3, the negation enters to the set of agents where p is true. Thus, we have

$$G(p \wedge \neg p) = G(p) \cap G(\neg p) = G(3) \neq \emptyset$$

Therefore, the contradiction is not empty. For set 1 we have

$$G(p \vee \neg p) = G(p) \cup G(\neg p) = G(2,3,4) \subseteq \text{Universe}$$

$$G(1) = G(2,3,4)^c$$

Hence, the tautology does not cover all sets of agents. If $C_1 \neq C_2$ then we may have an *ignorance* state which leads to the situation where tautology is not true. The agent can be *unaware* about presence of two different criteria, C_1 and C_2 .

Thus, we can have a *contradiction* state, $g(p, C_1) \neq g(p, C_2)$. Different aspects of concepts of Ignorance or contradiction states are discussed in [15]. Using a concept of the second order of conflict we can explain the contradiction of Zadeh's rule

$$\mu(p \wedge \neg p) = \min(\mu(p), \mu(\neg p)) = \min(\mu(p), 1 - \mu(p)) \geq 0$$

to the classical logic, where $\mu(p \wedge \neg p) = 0$.

Zadeh's rule produces $\mu(p \wedge \neg p) > 0$ for the contradiction $p \wedge \neg p$. In the AUT agent approach to uncertainty, the rule of min for AND operation is valid only when $G(\neg p) \subseteq G(p)$ that is

$$|G(p \wedge \neg p)| = |G(p) \cap G(\neg p)| = \min(|G(p)|, |G(\neg p)|) = |G(\neg p)|$$

$$\text{if } G(\neg p) \subseteq G(p) \text{ then } |G(p)| \geq |G(\neg p)|$$

All agents in $G(\neg p)$ are in a state of logical self-conflict. With the introduction of the logically self-conflicting agents or contradictory agents, we solve the Zadeh's incoherence among the conjunction and the negation operations.

10.6 Sugeno Negation and Negation with Agents at the Second Order of Conflict

For the Sugeno negation,

$$\mu(\neg p) = \frac{1 - \mu(p)}{1 + \lambda \mu(p)} \quad \text{where } \lambda = [-1, \infty]$$

we have

$$\mu(\neg p) + \mu(p) = \frac{1 - \mu(p)}{1 + \lambda \mu(p)} + \mu(p) = \frac{1 + \mu(p)^2}{1 + \lambda \mu(p)}$$

and

$$\mu(\neg p) + \mu(p) - 1 = \lambda \mu(p) \left(\frac{\mu(p) - 1}{1 + \lambda \mu(p)} \right)$$

Therefore,

$$\text{if } \lambda \geq 0 \text{ then } \mu(\neg p) + \mu(p) \leq 1$$

$$\text{if } \lambda < 0 \text{ then } \mu(\neg p) + \mu(p) > 1$$

Fig. 10.9 shows the sets of agents for two cases of the parameter λ

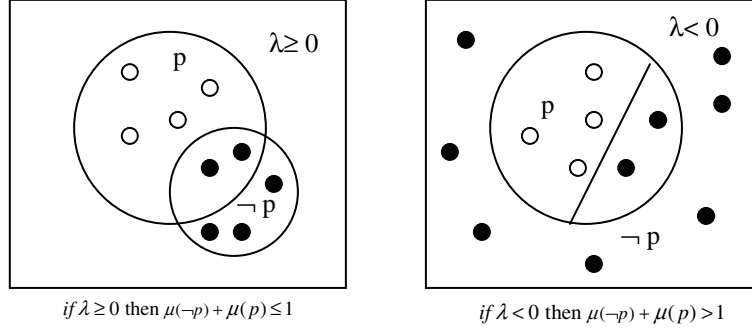


Fig. 10.9. Example for the Sugeno negation by set of agents

In Fig. 10.8 black circles show the agents for which $\neg p$ is true and white circles show agents for which p is true for the situations with $\lambda > 0$, and $\lambda < 0$, respectively.

10.7 Evidence Theory and Agents

In the evidence theory any subset $A \subseteq U$ of the universe U is associated with a value $m(A)$ called a basic assignment probability such that

$$\sum_{k=1}^{2^N} m(A_k) = 1$$

Respectively, the belief $Bel(\Omega)$ and plausibility $Pl(\Omega)$ measures for any set Ω are defined as

$$Bel(\Omega) = \sum_{A_k \subseteq \Omega} m(A_k), \quad Pl(\Omega) = \sum_{A_k \cap \Omega \neq \emptyset} m(A_k)$$

Thus, Bel measure includes all subsets that are inside Ω , The Pl measure includes all sets with non-empty intersection with Ω . In the evidence theory one and only one agent is associated with any element in the set Ω . Therefore, we have the situation shown in Fig. 10.10.

Fig. 10.10 shows set A that is overlapped with set Ω , which is divided in two parts: one inside Ω and another one outside. For the belief measure we exclude set A , thus we exclude the false state (f, f) and a logical self-conflicting state (t, f) . But for the plausibility measure we accept the (t, t) and the self-conflicting state (f, t) .

In the belief measure, we *exclude* any possible self-conflicting states, but in the plausibility measure, we *accept* self-conflicting states.

There are two different criteria to compute the belief and plausibility measures in the evidence theory. For belief C_1 criterion is related to set Ω , thus C_1 is true for the cases inside Ω . The second criterion C_2 is related to set A , thus C_2 is false for the cases inside A .

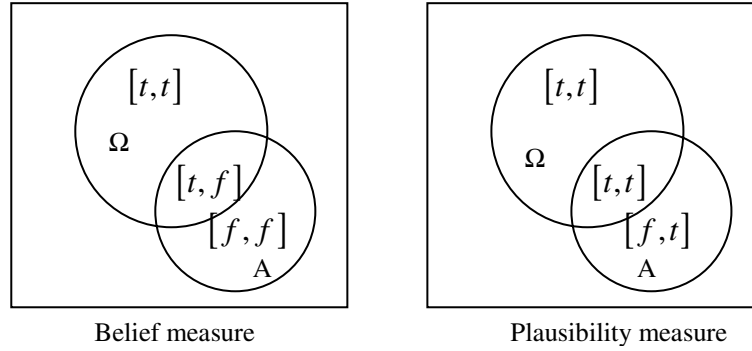


Fig. 10.10. Example of irrational agents or self-conflicting agent for the belief theory and plausible measures

Now we can put in evidence the logically self-conflicting state (t, f) , where we eliminate A also if it is inside Ω . The same is applied to the plausibility, where C_2 is true for cases inside A . In this situation, we accept a logically self-conflicting situation (f, t) , where cases in Ω are false but cases in A are true.

10.8 Rough Set Theory and Agents

Below we use the logical self-conflict to study the rough set. See Fig. 10.10. In Fig. 10.11, the rough set is defined by two criteria. The first criterion C_1 assumes that p is true if x is inside the partition of the space, the second criterion assumes that p is true if x is inside the set Ω .

The rough sets have internal and external measures. When elements of the partition are accepted without any self-conflict, we get an *internal measure*. When we accept the elements of the partition having the self-conflicting agents, we obtain the *external measure*.

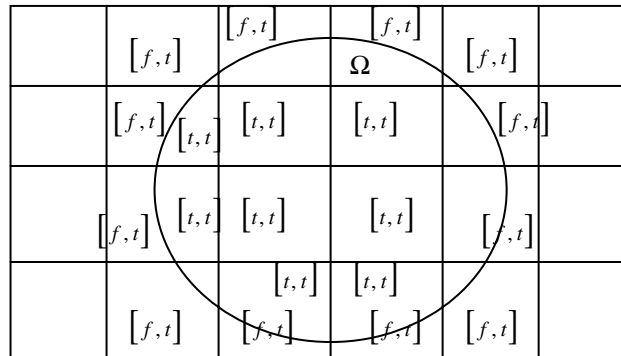


Fig. 10.11. Example of irrational agents for the Rough sets theory

10.9 AUT and Uniform Presentation of Uncertainty Theories

The agent approach to uncertainty permits to model different types of uncertainty due to: (1) ability of a set of agents to fuse conflicting logic evaluations, (2) ability of an individual agent to be self-conflicting, and (3) ability to use different negations relative to the ordinary logic evaluations. The contradiction can be true because a self-conflicting agent can evaluate p with one criterion and evaluate $\neg p$ with another criterion. Thus, p can be true for one criterion and $\neg p$ can be true for another criterion. This generates the internal irrational state. The same takes place for the tautology, p can be false for one criterion or $\neg p$ can be false for another criterion. Therefore, $p \vee \neg p$ can be false. Again, for the internal conflict we have another irrational situation for which tautology can be false.

Fusion of evaluations provided by agents and the introduction the internal structure represented by a set of criteria appear similar to the quantum interference or superposition for which conflicting positions or states can converge in one meta-position or state generated by the interference of the positions with different weights. Thus, we can say that in quantum mechanics we have a logic superposition for the same particle in different positions.

The second and higher levels of conflicts in AUT expand the concept of the agent to the concept of *self* that is an individual agent with the internal structure (presented by a set of criteria C) and/or a *fused agent* (superposition of agents, particles).

In analogy, the entangled particles create a meta-particle which itself is composed from different correlated particles. Entangled particles can be modeled by using a set of criteria C , one for any particle in the entanglement state. In other words, entanglement generates an irrational situation, where contradiction is true and tautology is false. This irrationality is represented by the instantaneous modification of the states of two particles without material communication among particles. This paradox was represented by the famous Bell inequality and by Einstein-Rosen-Podolsky or EPR paradox [26]. Here quantum mechanics violates the principle of realism for which any correlation is associated with a material communication with maximum velocity at the speed of light.

Table 10.1. Agents and uncertainty theories

<i>Theory</i>	<i>Component that uses rational agents</i>	<i>Component that uses irrational agent and self-conflict states</i>
Probability theory	Probability measure	No
Fuzzy logic	Membership function	Membership function
Intuitionistic fuzzy sets	Membership function	Membership function
Sugeno negation	Membership function	Membership function
Evidence theory	Belief measure	Plausibility measure
Rough Sets	Internal measure	External measure
Quantum mechanics		Particle superposition, entanglement

One of the important results of AUT is the possibility to model different types of uncertainties (probability, fuzzy sets, evidence theory evaluations, and rough set evaluations) under the same umbrella of the agents. Table 10.1 summarizes components of these theories that are modeled in AUT by rational agents and those that are modeled in AUT by irrational agents. Some components can be modeled by both types of agents. A more detailed classification of these components with different levels of conflicts and criteria using AUT is the subject of future research.

10.10 Applications

Potential application fields of AUT include verification of software agents, robotics, quantum computing, sensor networks, social modeling such as terrorism, stock market trading, marketing, web-based recommending systems that help a user to select a product, and many other fields. In all these areas agents demonstrate both rational and irrational behavior.

Different aspects of uncertainty have been considered in many types of Multi-Agent Systems (MAS). One of them is the problem of coalition formation when agents are *uncertain* about the types or *capabilities of their potential partners*. In [27] it was approached by using the probabilistic technique of a Bayesian reinforcement learning (Bayesian agents) for the situation when coalitions are formed and tasks undertaken repeatedly. Bayesian agents are rational, but their information about other agents is not complete. Here the list of types or capabilities of potential partners (other agents) are known to the agent, but it is not known which particular type is a type of the specific partner.

Rational agents operating under uncertainty also considered in [28]. These agents maximize the expected performance criterion in the form of monetary a utility payoff matrix or payoff game matrix that express preferences of agents over goods, states, or money that depend on wealth levels of agents. The wealth level of agents is considered as *private* information while the monetary payoff game matrix as *public* information that is available to each agent. The uncertainty is coming from the fact that private information is not known to other agents. Agents must negotiate to find a *stable state* using only the *public information*. The internal information about the agent that is taken into account in [28] is the *risk preference type* (risk averse, risk neutral and risk seeking). Note that from our viewpoint, even agents with risk neutral preferences can be rational or irrational, while risk averse and risk seeking agents can exhibit irrational behavior more often. This can dramatically change the negotiation result when agents attempt to reach a stable equilibrium state under uncertain games. Thus, we see that incorporation of the AUT concept of irrational agents can be beneficial in these agent negotiation tasks.

The paper [29] investigates the effect of different rationality assumptions on the performance of co-learning by multiple agents in multi-step games with close interactions between agents where an agent learns not only its own value function but also those of other agents. This paper shows complex pattern of effects resulting from different levels of rationality assumptions.

A multi-agent planning problem under temporal constraints and uncertainty of admissible operations is studied in [30]. The aim is to transform the system into one of the goal states satisfying all specified constraints with given a description of the current state of the system, a set of possible actions on the system and a description of its goal states. Here the uncertainty is associated with the set of admissible operations that are not known completely, but each agent is again rational.

The impact of uncertainty in agent-based approach in Market Engineering is considered in [31] and the uncertainty management framework for the multi-agent system is proposed in [32].

Below we outline one of the possible processes of applying the AUT:

- (1) Identify a set of agents $G=\{g\}$;
- (2) Identify a set of propositions $P=\{p\}$ that agents evaluate;
- (3) Collect data: $g(p)$ values for g and p in G and P . It can be a complete data set or a representative sample;
- (4) Identify contradictory agents and compute agent set contradictory index (see section 10.3.1);
- (5) Identify irrational agents as defined in section 10.3.1;
- (6) Compute individual agent irrationality index $\gamma(g,p)$ defined in section 10.3.1;
- (7) Identify a set of criteria $C=\{C\}$;
- (8) Collect data $g(p, C_k)$ for g, p and C_k in G, P and C for irrational agents. It can be a complete data set or a representative sample;
- (9) Identify irrational agents for a set of criteria C as defined in section 10.3.1;
- (10) Compute individual agent irrationality index $\gamma(g,p,C)$ defined in section 10.3.1;
- (11) Test if the criteria C can explain irrational behavior of agents by finding criteria C_i and C_j for p such that $(g(p,C_i)=\text{True} \wedge g(\neg p, C_j)=\text{True})$ or $(g(p,C_i)=\text{False} \wedge g(\neg p, C_j)=\text{False})$.
- (12) Test if AND or OR operations on propositions can resolve irrational behavior of agents. The example below illustrates this situation.

Consider an example of a car selection by customer from the viewpoint of car dealers. The dealers know that customers exhibit irrational behavior quite often. They may not be able to change such irrational behavior, but can build a marketing strategy based on irrational preferences. Say, it was discovered that the customer considers car C as better than car A in spite of his/her preference $A>B$ and $B>C$. Here we have a case of violation of the rational behavior (transitivity violation) by the agent. In this situation, the dealer can build an advertisement around car C not car A. On the other hand, if a dealer wants to modify preferences of the customer the dealer may add another product to be sold together with cars to make A preferable and to get rational preference of the agent. This new product can be the extended warranty for the car, which is crafted to be better for car A. The AUT vector logic with multiple criteria can model such situation by adding a new criterion. Having multiple agents and the need to construct a marketing strategy that will work for many customers, the AUT fusion (aggregation) of vector evaluations can be used. The AUT also helps to understand the reasons of the irrational behavior of the agents by the use of a set of criteria.

10.11 Conclusion

Agents and environment are major sources of uncertainty in the multi-agent systems. A set of all decision alternatives and states of the environment typically are unknown to the autonomous agent. These are fundamental *external* (to the agent) sources of uncertainty. The probability theory assumes that all alternatives are known in advance (e.g., six sides of the dice), but their occurrence is known only partially (probabilistically). Another source of uncertainty in MAS is the agent's own *internal* reasoning abilities. If the agent does not evaluate alternatives rationally and does not reason rationally, then this creates uncertainty of the agent's decisions and actions. The AUT addresses both internal and external uncertainty challenges in multi-agent systems.

The hierarchy of sets of agents relative to levels of conflicts, self-conflicts, and irrationality proposed in this paper is intended to provide a base for several areas of further studies. These studies include the foundation of probability theory, fuzzy logic, and other types of uncertainties, as well as real-world interpretation of these theories including quantum mechanics and quantum computing.

The major weakness in fuzzy logic is its ad hoc nature with operations that are not justified in advance (as it is done in the probability theory), but adjusted in many different ways (by using many different t-norms and t-conorms). The novelty of this paper is not in introduction of new formulas for initial membership function values and operations with them, but in creating a base for their *interpretation* by means of conflicting and irrational agent as an internal part of the theory. Physics provides plenty of positive examples of such interpretations that can be inspiration examples for building comprehensive logics of uncertainty for agents.

A society of agents always has a degree of conceptual conflicts and any agent can be self-conflicting. This affects agents' abilities to use resources for identifying goals and acting. In this paper, we established the fundamental logic structure that can be used for making logic decisions in a conflicting multi-agent environment. The new contributions of this paper consist of:

- 1) Introduced the agent uncertainty theory (AUT) to model the uncertainty in the society of agents;
- 3) Discovered connection among different types of uncertainties by using agents;
- 2) Discovered connection between classic logic and fuzzy logic using AUT that allows providing empirical, physical interpretations for them and distinction of their domains;
- 4) Built a new model of many-valued logic and uncertainty by using concepts of conflicts among agents and irrationality of agents;
- 5) Identified a way to model quantum mechanics by using agents and many-valued logic.

In the future research we plan to study more complex structures with more criteria and higher order of conflicts. In fact, a set of N criteria can generate $O(2^N)$ the logically self-conflicting states and very complex conflict situations. We expect that in near future this will be an active area of new fundamental research and discoveries.

References

1. Carnap, R., Jeffrey, R.: *Studies in Inductive Logics and Probability*, vol. 1, pp. 35–165. University of California Press, Berkeley (1971)
2. Fagin, R., Halpern, J.: Reasoning about Knowledge and Probability. *Journal of the ACM* 41(2), 340–367 (1994)
3. Edmonds, B.: Review of Reasoning about Rational Agents by Michael Wooldridge. *Journal of Artificial Societies and Social Simulation* 5(1) (2002), <http://jasss.soc.surrey.ac.uk/5/1/reviews/edmonds.html>
4. Ferber, J.: *Multi Agent Systems*. Addison Wesley, Reading (1999)
5. Gigerenzer, G., Selten, R.: *Bounded Rationality*. MIT Press, Cambridge (2002)
6. Halpern, J.: *Reasoning about uncertainty*. MIT Press, Cambridge (2005)
7. Hisdal, E.: *Logical Structures for Representation of Knowledge and Uncertainty*. Springer, Heidelberg (1998)
8. Resconi, G., Jain, L.: *Intelligent agents*. Springer, Heidelberg (2004)
9. Resconi, G., Klir, G.J., Clair, U.S.: Hierarchical uncertainty metatheory based upon modal logic. *International J.Gen.Systems* 21, 23–50 (1992)
10. Resconi, G., Klir, G.J., Clair, U.S., Harmanec, D.: In the integration of uncertainty theories. *Intern. J. Uncertainty Fuzziness knowledge-Based Systems* 1, 1–18 (1993)
11. Resconi, G., Klir, G.J., Harmanec, D., Clair, U.S.: Interpretation of various uncertainty theories using models of modal logic: a summary. *Fuzzy Sets and Systems* 80, 7–14 (1996)
12. Harmanec, D., Resconi, G., Klir, G.J., Pan, Y.: On the computation of uncertainty measure in Dempster-Shafer theory. *International Journal General System* 25(2), 153–163 (1995)
13. Resconi, G., Murai, T., Shimbo, M.: Field Theory and Modal Logic by Semantic field to make Uncertainty Emerge from Information. *Int. Journal General System* 29(5), 737–782 (2000)
14. Resconi, G., Turksen, I.B.: Canonical Forms of Fuzzy Truthhoods by Meta-Theory Based Upon Modal Logic. *Information Sciences* 131, 157–194 (2001)
15. Resconi, G., Kovalerchuk, B.: The Logic of Uncertainty with Irrational Agents. In: *Proc. of JCIS-2006 Advances in Intelligent Systems Research*. Atlantis Press, Taiwan (2006)
16. Kahneman, D.: Maps of Bounded Rationality: Psychology for Behavioral Economics. *The American Economic Review* 93(5), 1449–1475 (2003)
17. Kovalerchuk, B.: Analysis of Gaines' logic of uncertainty. In: Turksen, I.B. (ed.) *Proceeding of NAFIPS 1990, Toronto, Canada*, vol. 2, pp. 293–295 (1990)
18. Kovalerchuk, B.: Context spaces as necessary frames for correct approximate reasoning. *International Journal of General Systems* 25(1), 61–80 (1996)
19. Kovalerchuk, B., Vityaev, E.: *Data mining in finance: advances in relational and hybrid methods*. Kluwer, Dordrecht (2000)
20. Wooldridge, M.: *Reasoning about Rational Agents*. MIT Press, Cambridge (2000)
21. Montero, J., Gomez, D., Bustine, H.: On the relevance of some families of fuzzy sets. *Fuzzy Sets and Systems* 16, 2429–2442 (2007)
22. Atanassov, K.T.: *Intuitionistic Fuzzy Sets*. Springer, Heidelberg (1999)
23. Flament, C.: *Applications of graphs theory to group structure*. Prentice Hall, London (1963)
24. Ruspini, E.H.: A new approach to clustering. *Information and Control* 15, 22–32 (1999)
25. Priest, G., Tanaka, K.: Paraconsistent Logic. *Stanford Encyclopedia of Philosophy* (2004), <http://plato.stanford.edu/entries/logic-paraconsistent>

26. D’Espagnat, B.: *Conceptual Foundation of Quantum mechanics*, 2nd edn. Perseus Books (1999)
27. Chalkiadakis, G., Boutilier, C.: Sequential Decision Making in Repeated Coalition Formation under Uncertainty. In: Padgham, Parkes, Müller, Parsons (eds.) *Proc. of 7th Int. Conf. on Autonomous Agents and Multiagent Systems (AA-MAS 2008)*, Estoril, Portugal, May 12-16 (2008),
<http://eprints.ecs.soton.ac.uk/15174/1/BayesRLCF08.pdf>
28. Soo, V.-W.: Agent negotiation under uncertainty and risk. In: Soo, V.-W., Zhang, C. (eds.) *PRIMA 2000. LNCS*, vol. 1881, pp. 31–45. Springer, Heidelberg (2000)
29. Sun, R., Qi, D.: Rationality assumptions and optimality of co-learning. In: Soo, V.-W., Zhang, C. (eds.) *PRIMA 2000. LNCS*, vol. 1881, pp. 61–75. Springer, Heidelberg (2000)
30. Baki, B., Bouzid, M., Ligęza, A., Mouaddib, A.: A centralized planning technique with temporal constraints and uncertainty for multi-agent systems. *Journal of Experimental & Theoretical Artificial Intelligence* 18(3), 331–364 (2006)
31. van Dinther, C.: *Adaptive Bidding in Single-Sided Auctions under Uncertainty: An Agent-based Approach in Market Engineering (Whitestein Series in Software Agent Technologies and Autonomic Computing)*, Birkhäuser, Basel (2007)
32. Wu, W., Ekaette, E., Far, B.H.: Uncertainty Management Framework for Multi-Agent System. In: *Proceedings of ATS 2003* (2003),
<http://www.enel.ucalgary.ca/People/far/pub/papers/2003/ATS2003-06.pdf>
33. Colyvan, M.: The Philosophical Significance of Cox’s Theorem. *International Journal of Approximate Reasoning* 37(1), 71–85 (2004)
34. Colyvan, M.: Is Probability the Only Coherent Approach to Uncertainty? *Risk Analysis* 28, 645–652 (2008)