

Article

Self-Emotion-Mediated Exploration in Artificial Intelligence Mirrors: Findings from Cognitive Psychology

Gustavo Assuncao ^{1,2,*} , Miguel Castelo-Branco ^{1,3}  and Paulo Menezes ^{1,2} 

¹ Department of Electrical and Computer Engineering, University of Coimbra, 3030-790 Coimbra, Portugal; mcbranco@fmed.uc.pt (M.C.-B.); pm@deec.uc.pt (P.M.)

² Institute of Systems and Robotics (ISR), DEEC, 3030-290 Coimbra, Portugal

³ Institute for Biomedical Imaging and Translational Research (CIBIT), ICNAS, 3000-548 Coimbra, Portugal

* Correspondence: gustavo.assuncao@isr.uc.pt

Abstract

Background: Exploration of the physical environment is an indispensable precursor to information acquisition and knowledge consolidation for living organisms. Yet, current artificial intelligence models lack these autonomy capabilities during training, hindering their adaptability. This work proposes a learning framework for artificial agents to obtain an intrinsic exploratory drive, based on epistemic and achievement emotions triggered during data observation. **Methods:** This study proposes a dual-module reinforcement framework, where data analysis scores dictate pride or surprise, in accordance with psychological studies on humans. A correlation between these states and exploration is then optimized for agents to meet their learning goals. **Results:** Causal relationships between states and exploration are demonstrated by the majority of agents. A 15.4% mean increase is noted for surprise, with a 2.8% mean decrease for pride. Resulting correlations of $\rho_{surprise} = 0.461$ and $\rho_{pride} = -0.237$ are obtained, mirroring previously reported human behavior. **Conclusions:** These findings lead to the conclusion that bio-inspiration for AI development can be of great use. This can incur benefits typically found in living beings, such as autonomy. Further, it empirically shows how AI methodologies can corroborate human behavioral findings, showcasing major interdisciplinary importance. Ramifications are discussed.

Keywords: exploration; artificial emotion; general artificial intelligence; reinforcement learning; intrinsic drives



Academic Editor: Sharifu Ura

Received: 1 July 2025

Revised: 8 August 2025

Accepted: 2 September 2025

Published: 9 September 2025

Citation: Assuncao, G.; Castelo-Branco, M.; Menezes, P. Self-Emotion-Mediated Exploration in Artificial Intelligence Mirrors: Findings from Cognitive Psychology. *AI* **2025**, *6*, 220. <https://doi.org/10.3390/ai6090220>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

Recent advances in AI have led to a surpassing of conventional methodologies and disruption of various human-powered fields. These are becoming more and more digital, from medical pathology [1] to industrial smart systems [2] and even socio-emotive companionship [3]. However, this prosperity is fickle as the existing AI methodology is largely ineffective when devoid of human guidance [4]. This is a clear indicator that research on AI learning should move closer to metalearning. Additionally, the parameter explicitness required of model designers should be considered when building training procedures for artificial agents (naturally without intrinsic motivation). In this context, we postulate artificial emotion as a missing catalyst of exploratory behavior in AI and develop primary work on how to use it for that goal. This particular characteristic can enable agents to focus on data in accordance with their needs/interests, given how appraisal and attention changes correlate to alter sensory processing of stimuli in an effect dubbed emotional

salience [5]. In fact, exploration is already known to be a fundamental aspect of cognitive development and independent behavior in human beings [6] given its contribution to the acquisition of knowledge. For instance, confirmation biases influence how information is sought to ratify prior beliefs and inference [7]. Consequently, for AI to grow autonomous, researchers should strive to develop methodologies congruent with biological processing and optimize informational search.

Scrutiny of epistemic/achievement states (i.e., emotions pertaining to the generation of knowledge and a sense of success) for the purpose of benefiting exploration tactics in AI has been attempted, though only from a few general perspectives. Approaches have applied model behavior differences as criteria to determine whether data is adequate for classification [8,9], and divergence in transition probability has been employed as learning reinforcement [10]. There are also instances of states being inferred from intrinsic reward [11] to drive curiosity in exploration. Results were interpreted as congruent with emotion in real life, where internal cognitive conditions related with emotion (e.g., incongruity and expectancy) do influence exploration [12,13]. However, these interpretations only posit conditions as causes of emotional variation, which then impacts exploration. Hence, emotion is only implicit and does not benefit from the better contextual adaptability, informational density, or socio-communicative relevance that explicitness can bring. This leads to low generalizability. Practical applications are lost, such as the impact over learning aspects besides exploration, human understanding of AI decision-making, or the ability to further integrate environmental information. Contrastingly, emotional influence is well acknowledged for several biological factors [14]. Thus, reproducing conditions for direct emotion manifestation and consequential impact may represent a better methodology than current state-of-the-art approaches.

This improvement is corroborated by works which have achieved greater performances from autonomous emotion-mediated parameter optimization. For example, works on the prediction error, learning rate, or actual reward have presented emotional modulation based on the difference between short- and long-term average of reward entropy over time [15] or emotion quantization as a linear combination of separate power levels [16]. A mere difference in visual stimuli has also been made to influence valence–arousal pairs [17,18], which then affect learning parameters. Moreover, the basing of emotion meta-learning techniques on neurophysiology, where researchers strive to replicate limbic circuitry and neuromodulation, has been shown to provide performance advantages in decision-making [3]. It is a type of strategy that benefits from interdisciplinarity, yet it is uncommon. As such, it could provide an edge when developing new parameter calculation techniques and contribute to emotion-mediated AI progress.

Linking the lack of learning autonomy in artificial agents with an observable influence of intrinsic emotional drives in living beings, there is motivation to emulate the latter in an attempt to mitigate the former. This emulation should build on knowledge already established by neuropsychological studies as a bootstrapping point. Furthermore, it should remain task-agnostic yet still provide some type of advantage for agent learning, either in terms of autonomy or efficiency. In order to tackle the challenges of endowing AI with emotion-mediated intrinsic driving and study its outcomes, this paper considers the following two research questions:

- RQ 1 : How can emotion be represented in artificial agents, and how will its influence over their behavior be evaluated so that results may be valid and comparable to human behavioral studies? To achieve this, we first build epistemic and achievement emotion functions based on links demonstrated in cognitive psychology. These are applied to a learning framework, whose behavior is evaluated under conditions similar to those of human studies [19–22].

- RQ 2: Will the manifested correlations between emotion and exploratory behavior be useful for data processing by agents, similar to what happens with human beings? To understand this we integrate the framework in a learning loop, replicated over a large number of agents, and assess emergent correlations. Parallelism is then drawn between the former and reported human behavior.

Answering these questions is meant to contribute a valid technique for artificial agents to explore data more autonomously. It is also meant to spark more interest in human–AI interdisciplinary studies. The rest of the paper is organized as follows: Section 2 briefly introduces the topic of emotion–behavior studies in psychology after framing our approach within the context of AI. Section 3 describes the design of each framework component, subsequently overviews the experimental arrangement inspired by human studies. Section 4 details the experimental procedure and obtained results, followed by Section 5, which discusses them. Finally, Section 6 concludes the paper.

2. Background and Related Work

This section overviews the neuropsychological motivation behind our approach and its importance for interdisciplinary analogy. It further presents works tackling learning mediation by emotion in artificial intelligence, framing its usefulness and exposing the novelty of this work.

2.1. From Psychology to AI

The study of links connecting epistemic and achievement emotions (i.e., states pertaining to the generation of knowledge and a personal sense of success) to exploratory behavior has been an active topic of research in cognitive psychology [19–22]. In behavioral studies, epistemic emotions are commonly observed transitioning into confused and curious demeanors [23]. They also potentially lead to heightened motivation and the pursuit of success [24]. Depending on its outcome, the pursuit can lead to pride or shame. If confronted with information contradictory of internalized knowledge (i.e., high-confidence errors), people can manifest surprise supplanted by unexpected error outcomes. All in all, experiencing a positive outcome in a task predisposes humans into seeking similar internal reactions and corresponding scenarios. This exploratory behavior happens for both epistemic and achievement cognitive paths [20], albeit with potentially different objectives. The process is regulated by the brain’s reward system, which adapts its signaling proportionally to the aforementioned conditions [25]. Considering those are reproducible in deep learning, epistemic and achievement emotions may be integrated in AI for exploratory benefits.

For studying the effects of emotion over human behavior, the methodology is straightforward. It typically relies on tasks [26], such as classification and trivia, designed to induce relevant scenarios. Specifically in the works by Vogl et al. [20,21], adulteration of common knowledge statements is presented in a veracity assessment task. These statements induce errors and trigger epistemic/achievement states, which translate into exploration. For instance, uninformed human participants being presented incorrect statements became surprised, as confidence on their personal knowledge clashed with mistakes. Complementarily, correct responses prompted a sense of pride. Either scenario sparked exploration, as demonstrated by requests for additional information.

AI-oriented tasks are evidently different from those humans perform in experiments. Regardless, this does not invalidate adaptation to a machine friendlier format, so comparable observations may be taken [27]. This is corroborated by the fact that AI is considered a valid framework on which to explore a range of neuropsychologic phenomena [28,29]. Our contribution adds to this by adapting Vogl et al. experimental conditions for an AI application. Confident participants can be represented by models with near-perfect per-

formance. Cognitive tasks can be directly emulated, for example, as simple classifications. Naturally, so can knowledge adulteration, which enables emotional triggering. With this notion, the problem became a matter of adequately integrating emotion in artificial agents. Then these can be validly regarded as participants in a cognitive psychology experiment.

2.2. *Emotion-Driven Learning*

Most studies integrating emotion in AI rarely draw from robust psychological and neurophysiological foundations. Fewer employ it as a driver of learning. Still, our work relates with some studies. For instance, in [30], the authors developed a latent dynamics model by calculating the dissimilarity from its posterior to prior beliefs. This rewarded exploration when it occurred. Schillaci et al. also presented a process for estimating the change in prediction error (PE) as a metric of learning progress [31]. This enabled agents to shift attention towards more interesting (change-inducing) goals when progress was inadequate. In [32], the authors suggest a form of intrinsic rewarding reliant on competence progress, which is analogous to achievement gratification. Based on it, agents can decide how to explore their goal space. While these approaches advance autonomous exploration, they overlook its intrinsic driving background, such as emotion. This omission limits contextual relevance and narrows the applicability of the resulting metrics [3]. They also disregard most parallelism with living systems, despite their objective usefulness for understanding exploratory origins. Besides potentially slowing AI advancement, this also incapacitates interdisciplinary understanding, since inspiration from biological functioning can help postulate novel theories on how cognition develops from basic neural activity [33].

Our contribution is two-fold, as the work bears considerable interdisciplinary interest besides being advantageous towards autonomous AI. First, we propose building agents in accordance with neuronal structuring. This novelism can mediate exploration, without the hindrances of related work. It also increases the plausibility that resulting exploratory tendencies are useful for understanding real neuronal arrangements [34]. Second, we further innovate by adapting the experimental conditions under which psychological studies assess emotion–exploration relationships in humans for AI agents. This is more akin to the trend of [35], where psychological findings were considered as a basis for designing AI experiments. Despite the necessary adaptations for a standard AI methodology, the resulting framework can adequately demonstrate that autonomous exploration is possible. Moreover, it can appropriately corroborate psychological findings and provide a basis for other hypotheses to be considered, which are otherwise not easily achieved in behavioral studies.

2.3. *Ethics*

Integrating emotion in AI systems introduces a layer of behavioral modulation that aligns with human interaction. This can have several benefits, as specified, but also raises important ethical concerns. One such concern is unpredictable behavior. As decision-making ties with emotional metrics, optimization of underlying correlations may deviate from what is safe into what agents need for success. In a controlled experimental environment, this is unlikely to cause harm. However, in the real world, there may be risks to human–agent interactions and potential boundary violation, social or physical.

The anthropomorphization of AI with emotion can also blur distinctions between real and simulated feeling. This can entail ascription of trust and compliance with systems liable to external manipulation. This can challenge broader social norms around responsibility and accountability, as it becomes unclear whether the system, designer, or improper manipulator should be held accountable for decision outcomes. While ethical concerns are valid, emotion-driven AI offers clear benefits when used responsibly. Emotions can serve as functional signals that enhance adaptability, learning, and human interactions, provided

systems are designed with transparency. For an in-depth overview of AI-related ethics, readers are invited to check [36].

In summary, our paradigm was expected to demonstrate a causal relationship where epistemic or achievement emotions served as mediators of exploratory behavior, mimicking the findings reported by Vogl et al. [20]. Novelism stems from how the proposed framework is built, as well as from the comparison of emergent behaviors with those of humans. Finally, the resulting impact over agent knowledge acquisition and overall behavior was also considered useful contributions for AI autonomy.

3. Materials and Methods

The proposed framework consists of a task-oriented module, whose rate of exploration is dictated by an actor–critic module. The latter derives this rate from performance-based emotional scoring. The following sections overview each component of the system, namely how the emotional functions were replicated from psychology observations, what composes the framework, and how the learning cycle is designed. Figure 1 displays the proposed framework as a reference.

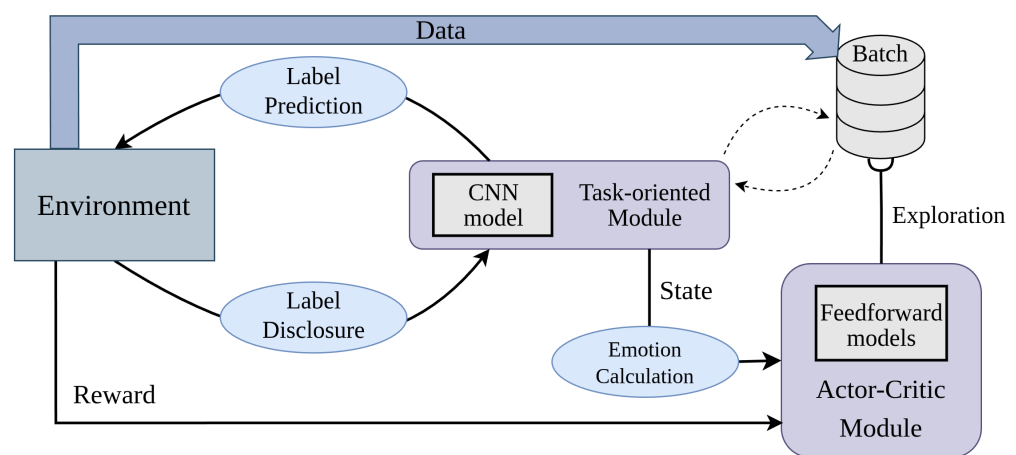


Figure 1. Framework overview, demonstrating how its components interact with the environment and compute an exploratory ration from performance-based emotion.

3.1. Replication of Surprise and Pride

We propose deriving epistemic and achievement emotion from the standard performance metrics of a deep learning methodology, interpreting them as underlying cognitive conditions. Specifically, testing accuracy reflects the adequacy of a model towards some task by gauging overall correctness over unseen data. Therefore, it may be employed as a pointer of error and achievement. In this case, escalation in accuracy scores can be interpreted as increasing success. Contrastingly, de-escalation entails a decrease in success. Variations in the feeling of pride should therefore match variations in accuracy, corresponding to personal achievement or lack thereof [37]. This accuracy–pride match can entail a curve with a positive slope and unknown convexity, with small variations. Factoring in confidence, besides accuracy, can broaden the set of representable emotions. For instance, high-confidence errors trigger the feeling of surprise, as derived from the cognitive incongruity explained previously. Additionally, insecure attainment of success may also induce surprise [38]. In these scenarios, a respective decrease or increase in confidence will instead lower surprise. A saddle-like behavior can therefore describe this emotion, as polarized variations of accuracy and confidence together imply intense values of surprise, whilst matching magnitudes of the two indicate a reduction or lack of this emotion. This view regarding surprise and pride is widely backed by the cognitive psychology literature [12,13,20,21].

Several different functions can represent the behavior described for surprise and pride. Here, a single set of functions that fulfill the requirements was selected arbitrarily for the main experiments. Examples are shown in Figure 2 for reference. As described, these examples factored in performance metrics of the task-oriented module, which is inspired by the impact of action outcomes over emotion [39]. This partially answers RQ1, contributing a novel way for emotion to occur in artificial agents, making it both objective and integrable in other AI pipelines.

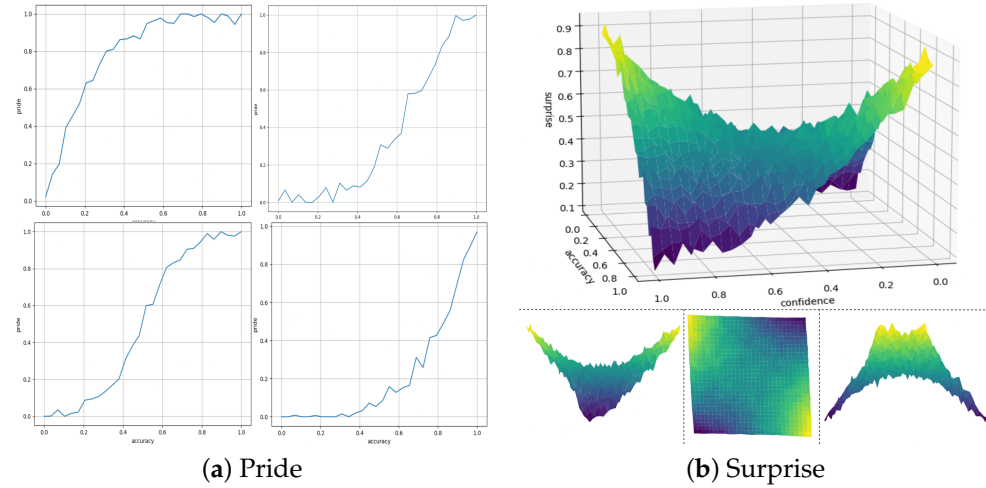


Figure 2. Example artificial emotion curves inspired by cognitive psychology [20,21]. (a) Positive pride slope based on increasing accuracy. (b) Saddle-like behavior of surprise based on accuracy and confidence.

3.1.1. Pride

Since accuracy has been shown to positively predict pride [20], making it is easy to compute and already fixed within $[0, 1]$, it can adequately model this emotion. To capture how pride might vary with increasing accuracy a , we propose a Gaussian-like bump function $P(a)$. This displays an overall upward trend, peaking near $a = 1$ (perfect accuracy), to reflect strong pride near high performance. We also include minor fluctuations to account for individual variability (e.g., personality and context). An example of such a function is

$$P: [0, 1] \rightarrow [0, 1]$$

$$a \mapsto \text{Clip} \left[(100 \cdot C_1)^{-(a-1)^2} + \mathcal{N}(\mu, \sigma^2) \right] \quad (1)$$

In the equation, $C_1 > 1$ controls how sharply pride rises as accuracy improves. The Gaussian noise $\mathcal{N}(\mu, \sigma^2)$ introduces individual variability. The clip function keeps pride within its natural bounds (0 and 1). This example is not meant to be precise or universal, but rather to illustrate a plausible emotional trajectory. Pride grows with success, but its exact path varies across individuals.

3.1.2. Surprise

Surprise depends on confidence, besides accuracy. Prior work shows that high-confidence errors are strong predictors of surprise [20]. To reflect this, we introduce a confidence score $c \in [0.8, 1]$, capturing the elevated confidence distribution yielded by the task-oriented model. Our surprise function $S(c, a)$ is proposed to provide high values when confidence and accuracy disagree (saddle-like behavior). This happens when an agent is

confident but wrong (high c and low a) or unsure but correct (low c and high a). It takes the following form:

$$S: [0,1]^2 \rightarrow [0,1]$$

$$c, a \mapsto \text{Clip} \left[\mathcal{T} \left(\mathcal{R} \left(a^2 - c^2 \right) \right) + 0.5 + \mathcal{N}(\mu, \sigma^2) \right] \quad (2)$$

Here, the core term $a^2 - c^2$ captures the disagreement between accuracy and confidence. A rotation $\mathcal{R}$ (by $45^\circ \pm C_2$) and translation $\mathcal{T}$ are applied to center and orient the surface so that surprise is maximized when accuracy and confidence diverge. As before, Gaussian noise introduces variability, and clipping keeps outputs within the range $[0, 1]$. The 0.5 term re-centers the output toward the range middle, reducing clipping distortion.

3.2. Framework Overview

In the framework, three models were implemented: the task-oriented, actor, and critic models. This combination was loosely inspired by human neural functioning during a psychological experiment, where part of the brain is focused on task success (e.g., prefrontal cortex), which is mediated by feedback from other regions (e.g., basal ganglia). It constituted each artificial participant, being fed data for classification with potentially wrong labels. This biological parallelism was purposeful so that the obtained results could be contrasted with behavioral studies with adequate validity, in line with answering RQ1 and fomenting human–AI interdisciplinarity in state-of-the-art research. Furthermore, this clear separation between the task-oriented module and emotion to exploration optimization expands on current AI exploration techniques, which usually combine them and disregard bio-inspiration.

3.2.1. Task-Oriented Module

This module is meant to carry out a cognitive task. Here, handwritten digit recognition was performed using the MNIST dataset [40] for the sake of simplicity. Other tasks would also be possible, as the module is task-agnostic. Convolutional and feedforward branches were combined in a VGG-like architecture [41] to first enhance visual cues representative of the content in images and then classify them as digits. The architecture is simple, with the layer arrangement outlined by Table 1.

Table 1. Task-oriented model architecture.

Layer	Type	Kernel/Units	Activation	Output Shape
1	Conv2D	3×3 (32 filters)	ReLU	$26 \times 26 \times 32$
2	MaxPooling	2×2	-	$13 \times 13 \times 32$
3	Conv2D	3×3 (64 filters)	ReLU	$11 \times 11 \times 64$
4	MaxPooling	2×2	-	$5 \times 5 \times 64$
5	Flatten	-	-	1600
6	Dropout	0.5	-	1600
7	Dense	128	ReLU	128
8	Dense	10	Softmax	10

In order to train this model, half of the MNIST training dataset was used unadulterated. Standard backpropagation was employed using the Adam optimizer [42] for 50 epochs, with a batch size of 64. Testing with the MNIST test set yielded over 99% accuracy and near-zero loss. The confidence distribution is heavily concentrated in the high-confidence bins, with over 98% of samples falling within the $[0.95, 1]$ range. This indicates that the model is highly confident in its predictions. Consequently, it is well-suited for applications

in our framework, where such consistent certainty parallels the behavior of a person that is highly confident as a result of repeated success. This further validates an answer to RQ1.

As for the second half of the MNIST dataset, it was used directly in the main surprise/pride experiments. This was adulterated so that 50% of its instances had random labels, making them different from the specific digits they represented. The adulteration is represented in Figure 3b. Evidently, this was performed at random indices, so there would be good sparseness of correct and incorrect labels. Hence, despite being technically correct when classifying any of the 30,000 images used in the experiment, the pre-trained task-oriented network would be met with disparate labels approximately 50% of the time. This discrepancy is meant as the emotional trigger in the framework.

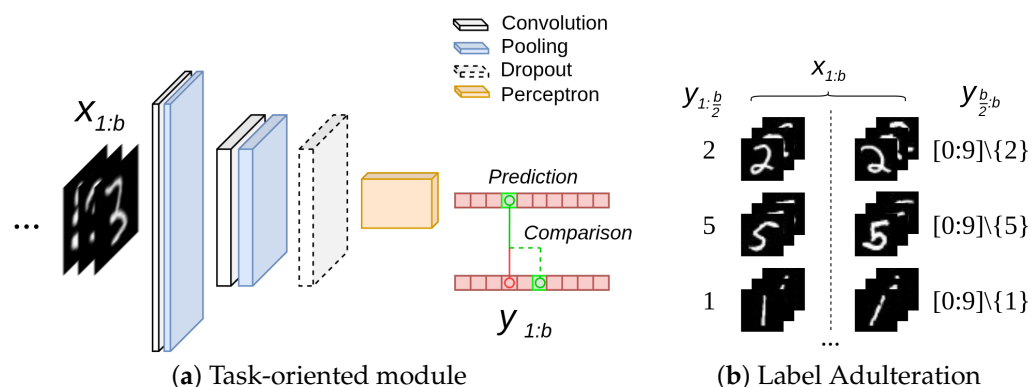


Figure 3. (a) The VGG-like model trained and employed in the framework for a classification task and (b) 50% label adulteration of the MNIST training data against respective visual content.

3.2.2. Actor–Critic Module

Our system requires a continuous evaluation of its own emotional state, processing it into an exploratory rate, as the most appropriate action. To achieve this, we drew inspiration from the habitual actor–critic dichotomy of the basal ganglia [43] to design decision-making neural modules in our artificial agents. Since we employ a form of directed exploration in our agent task, a deterministic approach would be more fitting than a stochastic one. Hence, deep deterministic policy gradients (DDPGs) [44] were implemented as AI parallels of the basal ganglia. This type of reinforcement learning (RL) methodology assumes separate networks: the critic model and the actor model. These collaborate to map the state to the action deterministically, attempting to maximize reward. Both the actor and the critic were implemented as multi-layer perceptrons focused on generating embeddings, which are then reduced, respectively, as an action or rectifying signal. These embeddings are parsed from the emotional state of the artificial agent, which is taken as the sole input for the actor, and tupled with the chosen action for the critic. An overview of this process is shown in Figure 4.

Within the RL paradigm, the agent state is activated as formula-driven surprise or pride scores. With it, the actor decides on an exploratory rate to be used by the task-oriented module. The critic then signals the actor regarding that rate's task-oriented usefulness. This translates as attentional shifting, as the actor is effectively warranting the task-oriented module to perform more/less intake of data in order to mitigate/potentiate emotional exacerbation. The resulting decision policy will therefore codify the exploratory rate in terms of epistemic or achievement emotion. We quantify this rate as a percentage to extract a fraction of the batch size, which is bound by $[0, Batch_{Max}]$. While not an exact representation of the basal ganglia and its related structures, this arrangement boasts similarities both architecturally and in terms of functioning of a human participant's decision-making process. A comparison with human study results becomes more valid, in line with RQ1.

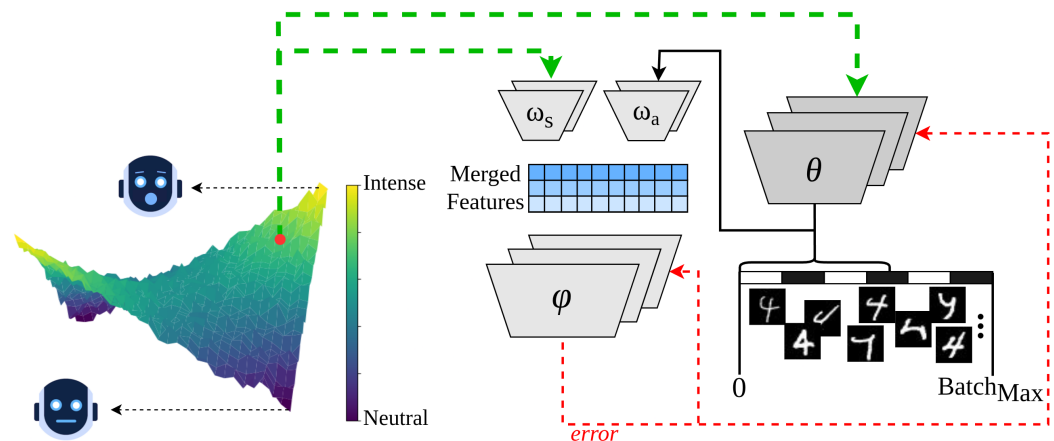


Figure 4. The actor–critic module. The actor θ computes an exploratory rate based on received emotional scoring. The critic $\omega - \phi$ scrutinizes this rate, generating feedback for itself and actor optimization towards task-oriented success.

3.3. Learning Cycle

Learning a correlation between surprise or pride and exploratory behavior involved all three models: the task-oriented, actor, and critic models. Since the first is already trained and highly confident in its predictions, its weights are frozen and used for forward passes only. To start, the set of MNIST with partially adulterated labels is made available to the task-oriented model for classification at instance-wise steps, which are mediated by the actor–critic module. This means an item is first picked randomly and processed to generate a corresponding label. The adulteration of labels ensures that a portion of the predictions is unavoidably incorrect. Still, confidence remains high due to the initial training process and weight freezing. The system is now able to experience both successful classification as well as incur high-confidence errors. Such circumstances induce emotional variation in accordance with the formulae described previously. This information is inputted to the actor, which will decide how much data of the same type should be analyzed subsequently (i.e., explored), using the output rate to compute a batch fraction. Processing of this fraction results in further emotional variation. Its comparison with single-instance states can yield insight into emotional progression during the cognitive task, in addition to its relationship with exploratory fluctuation overall.

The emotional scores of the system and the actor-derived exploratory rate are taken in by the critic for adequacy assessments. This process depends on whether the chosen rate improves the system's condition, contributing towards its objective. To specify this, we design a reward signal following the standard assumption that participants typically intend to perform well in the activities they perform, maximizing success. Therefore, reward varies analogously to common human functioning wherein reward-coding neurons respond to success from profitable decision-making [45] and maintenance of a homeostatic balance [46]. In our framework, each RL agent obtains a basis reward value, whose polarity corresponds to that of the difference between explored batch accuracy and single-instance accuracy. This ensures that exploration is useful only if yielding improvements in terms of task performance. Additionally, agents are provided with a sparse reward matching the variation in epistemic/achievement emotion, which occurs during a step. This either minimizes surprise or maximizes pride, complying with the free energy principle, which illustrates a necessity of self-organizing agents to reduce uncertainty in future outcomes [47]. The reduction can stem from knowledge diversification, which is boosted by an exploratory increase, or from near-complete reliance on current knowledge, where exploration is largely avoided.

4. Experiments and Results

The experimental application of our framework attempted to remain as close to Vogl's study as possible [20]. Since variations in nature/nurture naturally influence emotion and decision-making [48], it is important to account for personality variability. Thus, a total of 250 artificial agents were created, with distinct emotional functions obtained through parameter variation and added noise. These were employed for surprise and pride experiments separately. For either emotion function, $\mathcal{N}(0, 0.03)$ was employed, along with random combinations of C_1 and C_2 to generate varied artificial agents with individual differences while still following the same grand pattern. The learning cycle was applied to each artificial agent, with a reset occurring every 20 steps to match Vogl's 20-statement procedure. This reset marked the beginning of each learning episode, with a total of 100 episodes solidifying the robustness of our observations on emotion–exploration. Additionally, the actor and critic models used the Adam optimizer during the cycle run, with learning rates of 0.001 and 0.002, respectively. A replay buffer was also implemented here to reduce the variance from temporal correlations. Target networks were also implemented to help regularize learning updates. Overall results are depicted in Figure 5. The associations between exploratory behavior and epistemic/achievement emotions are analogous to findings reported in the original cognitive psychology study we strived to emulate [20].



Figure 5. Results for both surprise and pride, mirroring similar findings in cognitive psychology. Leftmost column: Episodic mean of emotion differential between single sample and subsequent batch analysis steps across all implemented agents over the entire learning cycle. Middle column: Mean cumulative reward obtained by agents at each episode of the cycle. Rightmost column: Mean actor behavior at the end of the learning cycle, correlating surprise or pride with exploration.

4.1. Outcome Analysis

First, model convergence was required to ensure behaviors learned by the artificial agents were not random. This was achieved for either emotion, as is demonstrated by the increase in and plateauing of cumulative reward over time (Figure 5 middle column). Specifically for surprise (top row), the initial reward is restricted to $[-8.58, 10.77]$, peaks at $\max_r^s = 19.87$, and ends within the range of $[-4.37, 19.77]$, with an early-stage dip minimum of $\min_r^s = -13.64$. For pride, the initial reward is encompassed by $[-15.97, 7.02]$, within which $\min_r^p$ is the minimum. The final reward here varies at $[-7.46, 18.95]$, though the overall maximum value $\max_r^p = 19.05$ is achieved shortly before. The average cumulative reward across agents also increases for both emotions. This demonstrates stable yet slight growth for pride. Contrastingly, surprise entails a short depression

in earlier episodes followed by a steady increase later on. Overall, these trends indicate that agents successfully learned to correspond states to actions in a useful way.

As episodic cumulative reward validates the success of the learning cycle, observed emotional fluctuations over time are also legitimized. These are presented in the first column of Figure 5. Expectedly, the initial variation is well-balanced for both emotions, as the number of increases matches that of decreases in the first episode. However, different outcomes manifest as episodes progress:

- Surprise (top trend): On average, bursts become less frequent by the final episode ($\Delta s = 38.52\%$). As such, it seems a reduction or stasis is favored over increases. Stasis is also progressively preferential in the first 10 episodes of surprise, falling back as decreases become more prominent later in the cycle.
- Pride (bottom trend): Overall, the average emotion variation among agents is minor yet still favoring an upwards tendency. Decreases occur fewer times ($\Delta p = 5.90\%$) between the first and last episodes. This percentage is largely taken over by stasis scenarios, with the number of emotional increases remaining mostly unchanged throughout the cycle.

4.2. Observed Correlations

The third column of Figure 5 evidences the impact of emotion over exploratory behavior. In spite of the emotional differences imposed via parameter fluctuation, a substantial number of artificial agents manifested similar behaviors after the cycle. A causation effect is most evident for the surprise experiment (top), which is akin to the emotional variation results. Averaging the decision-making behavior of all 250 agents displayed a 15.4% increase in exploration in response to greater surprise. This trend resonates with the 217 agents that learned positive correlations, outshining the remainder 33 who displayed negative correlations. Regardless of their variability, all instances were monotonic and either displayed a considerable increase (positive) or a limited decrease (negative). In the pride experiment (bottom), instances were likewise monotonic. However, agents here demonstrated a deflating exploratory effect. Behaviors encompassed a large amount of positive weak correlations, with few yet ample negative correlations. Specifically, 222 agents displayed a slight exploratory increase with pride. The impact proved minimal, as represented by the mean behavior of the subset. A smaller set of 22 agents decreased exploration between 25% and 75% towards null. Contrastingly, this caused a substantial negative change in mean behavior. Moreover, it is aided by the six remaining agents who manifested a more restrained reduction. The final result is a modest effect of a 2.8% overall average decrease in exploration for increasing pride, despite several weak positive correlations. Here medians can provide better insight by more accurately demonstrating negative cases as few yet hefty outliers against the whole.

Relationship strength between exploration and pride/surprise can be further demonstrated by measuring how each data pair correlates throughout a cycle. Here, we employ Spearman's correlation coefficient ρ [49] at each episode in an agent's cycle. This assesses if observed monotonic relationships between emotion and exploration (Figure 5 third column) become increasingly robust over time. Resulting coefficients for either experiment, with per-episode averaging of all 250 agent sample pairs, are shown in Figure 6. These demonstrated considerable variability in the strength of the correlation between exploration and surprise/pride. To mitigate this effect, a sliding window encompassing 40 episodes was applied to smoothen the trends and clarify strength progression. Both cases display near-zero coefficients in earlier episodes. However, this mean $\bar{\rho}$ increases for surprise while decreasing for pride, resulting in $\rho_{surprise} = 0.461$ and $\rho_{pride} = -0.237$ by the end of the cycles. This evidences a moderate positive correlation for surprise and exploration.

For pride, there is instead a negative and much weaker association with this behavior. Succinctly, this weakness is congruent with the disparity of having positive cases 10 times more common than negatives ones while still obtaining a negative correlation.

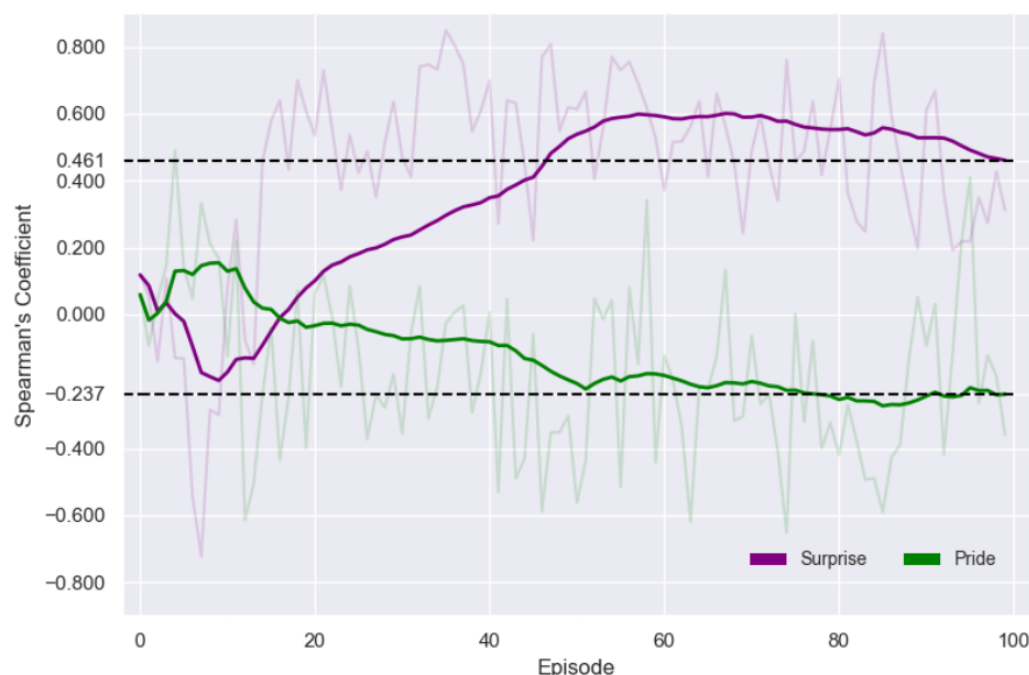


Figure 6. Agent episodic mean of Spearman's correlation coefficient between the actor-chosen exploratory rate and its causal surprise or pride score (pale), smoothed by a moving window of 40 samples (bold).

5. Discussion

Results hold a clearly substantial minimization of surprise. This conforms with the free energy principle, as it led to a strong direct correlation being learned between this emotion and exploration. Contrarily, maximization of pride was somewhat negligible. This was paralleled by the weak dampening relationship obtained for pride over exploration, with its strength stagnating closer to zero. Notwithstanding the latter, both experiments successfully produced artificial agents capable of self-mediating their exploratory behavior. They did so by exploiting internal emotional drives towards improved task performance.

The outcome of our experiment strongly resonated with reports on emotion-mediated human exploratory behavior. Specifically, Vogl's study [20] demonstrated a causation effect from surprise over exploration of knowledge. This was evidenced by the successively positive path coefficients obtained when assessing surprise to curiosity and curiosity to exploration effects. Additionally, the first and second versions of the study reported within-person correlation coefficients of 0.285 and 0.262 for this emotion and exploration. Even higher values were reported when considering curiosity intermediately. Relating these findings with our work, the 15.4% mean exploratory increase leveraged by the agents over growing surprise validates Vogl's postulated path relationships. Our Spearman's correlation results further support this, given the additional value proximity. For surprise, as the non-windowed coefficient mean across agents reaches 0.311 by the final episode, it becomes considerably close to the within-person correlation value of the first study (with the most participants). It also goes reasonably near that of the second study.

In terms of pride, Vogl's observations were also validated by our AI experiment. While negative correlation coefficients of -0.073 and -0.177 were reported in the first and second studies, respectively, these near-zero values indicate a weaker correlation between pride and exploration, if any. Though contradictory, path coefficients first indicated a weak but

positive causation effect followed by a stronger dampening impact later. This suggested that pride has a faint influence over exploration. Our results are congruent, as Spearman's correlation coefficient took longer to deviate from null compared to the surprise experiment, stagnating at a lower absolute value. Additionally, our 2.8% mean exploratory decrease over growing pride supports the higher likelihood of this emotion's dampening effect over exploration. While this was due to a smaller amount of negative strong correlations, the insignificance of a considerably larger amount of positive relationships aligns with pride's influence over exploration being modest. It could also be postulated that exploratory decrement is typical mostly during surges of pride. This is because, as a solo positive emotion, it may be damaging for cognitive performance [50], of which exploratory behavior is a key aspect [51]. Additional experimentation would be required to test such hypotheses. Regardless, under equivalent test conditions, our agents recognized a benefit in adopting behaviors akin to those exhibited by humans. These observations respond to RQ2, suggesting that emotion–exploration correlations manifested by agents are not only observable but also functionally identical. It is therefore plausible that these artificial correlations mirror their role in human cognition, being useful for data processing at least within comparable scenarios.

On a related note, Vogl's appealed for conceptual replication of their results to bolster generalizability [20,21]. That is what inspired RQ1, to which our proposed framework responds and meets the appeal. Also, employing an image classification task is evidently different from the general knowledge trivia scenarios devised by those authors. Thus, our results further support the conclusion that observed emotion–exploration correlations were not merely triggered by biased input stimuli. Moreover, cognitive incongruity was induced differently, as half of the data instances were assigned randomly incorrect labels at reset, rather than always being correct/incorrect. This meant contradictory information was a possibility, as samples with visually distinct content could be assigned the same label. Vogl et al. also stressed the importance of considering various indicators of knowledge exploration when attesting the validity of an epistemic/achievement source. Unlike deep learning, where exploration may be parameterized, psychology relies on observations of behavior to assess how exploration occurs and to what extent [19,20,38]. As our models objectively derive exploration from emotional scoring, this multi-indicator requisite was irrelevant in our approach, with no impact over findings. Finally, implementing artificial agents as participants in a within-“person” experiment is fundamentally distinct from human trialing. This constitutes further variability in comparison with Vogl's original study.

Limitations

This work bears limitations, given the variability of its many factors and the aim of interdisciplinary comparisons. For instance, while authors strived to propose general solutions for emotional representation, the functions may still introduce a potential bias which hinders generalizability. The introduction of noise can help reduce this bias, yet a more robust approach would be to consider other functions that also match the psychological descriptions of surprise and pride and to extend the experiment. A similar limitation is related with the usage of MNIST and classification as the cognitive task. We used a simplistic example, where high success rates are well-established in deep learning. Different outcomes could occur if considering a different task or using a more complex dataset (e.g., CIFAR), and thus bias may also exist. To mitigate these limitations, we intend to vary the emotional functions further and employ different datasets for classification so that more generalizable conclusions can be made in future work. Still on the topic of validation, our study is limited to 250 agents. Naturally, experiment expansion can also

entail a greater amount of agents in order to better reflect the diverse nature of living beings and their behaviors.

It is also important to note that artificial agents are unlikely to bear conscious metacognition and be refrained to supporting humans in complex tasks in the foreseeable future [27]. Thus, a comparison of results obtained from biological participants and AI should be approached with care. Given this limitation, we opt for the corroborating usefulness of agent observation. This can provide further scrutiny for hypotheses on human cognition and behavioral traits [52]. Finally, our system addresses each emotion individually as a precursor of exploration, whereas emotional states are most typically overlapping and combine effects over general behavior [3]. In the future, we intend to expand the emotional basis of exploration to account for that overlap, further benefiting from and validating cognitive psychology hypotheses. For instance, epistemic and achievement states could be employed in tandem, equitably, or via weighted contributions as inputs to an actor module. We speculate that emotional overlap is needed to stabilize correlations, causing more well-defined behaviors, similar to living beings. If observed, this could provide further insights into human–AI parallelism and serve as inspiration for cognitive psychology expansion.

6. Conclusions

This work focused on developing a deep learning framework for emotional decision-making over exploration. The architectural design was inspired by basal ganglia circuitry, with the foundation for emotional operation stemming from cognitive psychology. Emulation of human experimental conditions from psychological studies was conducted via an original learning cycle. This was applied to a novel deep learning framework, which was replicated several times for generalizability purposes. This proposal adequately solves RQ1, though others may still be possible. Furthermore, AI-learned correlations between epistemic/achievement states and exploration demonstrated close proximity with observations taken from human studies. Hence, for RQ2 we speculate the correlations to indeed be useful for AI agents, much like they are to their human counterparts. Our work additionally supports emotion-mediated learning, given its benefits to explainable AI and its autonomy. These include, for instance, greater contextual adaptability for agents to self-adapt beyond just exploration or a human-like understanding of AI decision-making.

In terms of future research, our framework follows the outlined benefits and others, since it can be adapted for studying other behavioral traits and their relationship with emotional drives. Exploitation or engagement may be mediated through variable emotion during cognitive operation. This can prove beneficial, similarly to how agents here learned to explore when it is seemingly more useful for their own reward objective. Finally, we speculate that further research on these topics can push AI closer towards full autonomy and general intelligence.

Author Contributions: Conceptualization, G.A., M.C.-B. and P.M.; methodology, G.A., M.C.-B. and P.M.; software, G.A.; validation, G.A.; data curation, G.A.; writing—original draft preparation, G.A.; writing—review and editing, M.C.-B. and P.M.; supervision, P.M. All authors have read and agreed to the published version of the manuscript.

Funding: This work has been financed by the PRR—Recovery and Resilience Plan—and by the Next-Generation EU European Funds, following NOTICE No. 02/C05-i01/2022, Component 5—Capitalization and Business Innovation (Mobilizing Agendas for Business Innovation)—under the project Greenauto (PPS10/PPS12/PPS13 with the reference 7255, C629367795-00464440).

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: Publicly available datasets were analyzed in this study. The full MNIST dataset can be found in <https://www.kaggle.com/datasets/hojjatk/mnist-dataset>, accessed on 11 January 2023. No new data were created in this study.

Acknowledgments: This work was partially supported by FCT under grant 2020.05620.BD and OE (National funds of FCT/MCTES) under project UIDP/00048/2020.

Conflicts of Interest: The authors declare no conflict of interest. The funders had no role in the design of the study; in the collection, analyses, or interpretation of data; in the writing of the manuscript, or in the decision to publish the results.

Abbreviations

The following abbreviations are used in this manuscript:

AI	Artificial Intelligence
CNN	Convolutional neural network
DDPGs	Deep Deterministic Policy Gradients

References

1. Baxi, V.; Edwards, R.; Montalto, M.; Saha, S. Digital pathology and artificial intelligence in translational medicine and clinical practice. *Mod. Pathol.* **2021**, *35*, 23–32. [\[CrossRef\]](#) [\[PubMed\]](#)
2. Sarker, I.H. AI-Based Modeling: Techniques, Applications and Research Issues Towards Automation, Intelligent and Smart Systems. *SN Comput. Sci.* **2022**, *3*, 158. [\[CrossRef\]](#) [\[PubMed\]](#)
3. Assuncao, G.; Patrao, B.; Castelo-Branco, M.; Menezes, P. An Overview of Emotion in Artificial Intelligence. *IEEE Trans. Artif. Intell.* **2022**, *3*, 867–886. [\[CrossRef\]](#)
4. Mühlhoff, R. Human-aided artificial intelligence: Or, how to run large computations in human brains? Toward a media sociology of machine learning. *New Media Soc.* **2019**, *22*, 1868–1884. [\[CrossRef\]](#)
5. Pauli, W.M.; Röder, B. Emotional salience changes the focus of spatial attention. *Brain Res.* **2008**, *1214*, 94–104. [\[CrossRef\]](#)
6. Nunnally, J.C.; Lemond, L.C. Exploratory Behavior and Human Development. In *Advances in Child Development and Behavior Volume 8*; Elsevier: Amsterdam, The Netherlands, 1974; pp. 59–109. [\[CrossRef\]](#)
7. Kaanders, P.; Sepulveda, P.; Folke, T.; Ortoleva, P.; Martino, B.D. Humans actively sample evidence to support prior beliefs. *eLife* **2022**, *11*, e71768. [\[CrossRef\]](#)
8. Kim, J.; Feldt, R.; Yoo, S. Guiding Deep Learning System Testing Using Surprise Adequacy. In Proceedings of the 2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE), Montreal, QC, Canada, 25–31 May 2019; pp. 1039–1049. [\[CrossRef\]](#)
9. Weiss, M.; Chakraborty, R.; Tonella, P. A Review and Refinement of Surprise Adequacy. In Proceedings of the 2021 IEEE/ACM Third International Workshop on Deep Learning for Testing and Testing for Deep Learning (DeepTest), Madrid, Spain, 1 June 2021; IEEE: New York, NY, USA, 2021; pp. 17–24. [\[CrossRef\]](#)
10. Achiam, J.; Sastry, S.S. Surprise-Based Intrinsic Motivation for Deep Reinforcement Learning. *arXiv* **2017**, arXiv:1703.01732.
11. Yin, H.; Chen, J.; Pan, S.J.; Tschitschek, S. Sequential Generative Exploration Model for Partially Observable Reinforcement Learning. *Proc. AAAI Conf. Artif. Intell.* **2021**, *35*, 10700–10708. [\[CrossRef\]](#)
12. Marshall, M.; Brown, J. Emotional reactions to achievement outcomes: Is it really best to expect the worst? *Cogn. Emot.* **2006**, *20*, 43–63. [\[CrossRef\]](#)
13. Pekrun, R. Achievement emotions: A control-value theory perspective. In *Emotions in Late Modernity*; Patulny, R., Bellocchi, A., Olson, R.E., Khorana, S., McKenzie, J., Peterie, M., Eds.; Routledge: Abingdon, UK, 2019; pp. 142–157. [\[CrossRef\]](#)
14. Damasio, A.R. *Descartes' Error: Emotion, Reason, and the Human Brain*; Avon: New York, NY, USA, 1994.
15. Huang, X.; Wu, W.; Qiao, H.; Ji, Y. Brain-Inspired Motion Learning in Recurrent Neural Network with Emotion Modulation. *IEEE Trans. Cogn. Dev. Syst.* **2018**, *10*, 1153–1164. [\[CrossRef\]](#)
16. Wang, C.; Mei, S.; Yu, H.; Cheng, S.; Du, L.; Yang, P. Unintentional Islanding Transition Control Strategy for Three-/Single-Phase Multimicrogrids Based on Artificial Emotional Reinforcement Learning. *IEEE Syst. J.* **2021**, *15*, 5464–5475. [\[CrossRef\]](#)
17. Hieida, C.; Horii, T.; Nagai, T. Deep Emotion: A Computational Model of Emotion Using Deep Neural Networks. *arXiv* **2018**, arXiv:1808.08447. [\[CrossRef\]](#)
18. Hieida, C.; Horii, T.; Nagai, T. Decision-Making in Emotion Model. In Proceedings of the Companion of the 2018 ACM/IEEE International Conference on Human-Robot Interaction, Chicago, IL, USA, 5–8 March 2018. [\[CrossRef\]](#)
19. Chevrier, M.; Muis, K.R.; Trevors, G.J.; Pekrun, R.; Sinatra, G.M. Exploring the antecedents and consequences of epistemic emotions. *Learn. Instr.* **2019**, *63*, 101209. [\[CrossRef\]](#)

20. Vogl, E.; Pekrun, R.; Murayama, K.; Loderer, K.; Schubert, S. Surprise, Curiosity, and Confusion Promote Knowledge Exploration: Evidence for Robust Effects of Epistemic Emotions. *Front. Psychol.* **2019**, *10*, 2474. [\[CrossRef\]](#) [\[PubMed\]](#)
21. Vogl, E.; Pekrun, R.; Murayama, K.; Loderer, K. Surprised–curious–confused: Epistemic emotions and knowledge exploration. *Emotion* **2020**, *20*, 625–641. [\[CrossRef\]](#)
22. Fitneva, S.A.; Slinger, M. Looking for a second opinion: Epistemic emotions and the exploration of information sources. *Proc. Annu. Meet. Cogn. Sci. Soc.* **2022**, *44*, 3950.
23. Muis, K.R.; Chevrier, M.; Singh, C.A. The Role of Epistemic Emotions in Personal Epistemology and Self-Regulated Learning. *Educ. Psychol.* **2018**, *53*, 165–184. [\[CrossRef\]](#)
24. Sznycer, D.; Cohen, A.S. How pride works. *Evol. Hum. Sci.* **2021**, *3*, e10. [\[CrossRef\]](#)
25. Schultz, W. Neuronal Reward and Decision Signals: From Theories to Data. *Physiol. Rev.* **2015**, *95*, 853–951. [\[CrossRef\]](#)
26. Hackman, J.R. Toward understanding the role of tasks in behavioral research. *Acta Psychol.* **1969**, *31*, 97–128. [\[CrossRef\]](#)
27. Korteling, J.E.H.; van de Boer-Visschedijk, G.C.; Blankendaal, R.A.M.; Boonekamp, R.C.; Eikelboom, A.R. Human-versus Artificial Intelligence. *Front. Artif. Intell.* **2021**, *4*, 622364. [\[CrossRef\]](#)
28. Dawson, M.R.W. Book Review—Computational neuroscience and cognitive modelling: A student’s introduction to methods and procedures, By Britt Anderson, Thousand Oaks, CA: Sage Publications, 2014. *Br. J. Psychol.* **2014**, *105*, 436–438. [\[CrossRef\]](#)
29. Dawson, M.; Dupuis, B.; Spetch, M.; Kelly, D. Simple Artificial Neural Networks That Match Probability and Exploit and Explore When Confronting a Multiarmed Bandit. *IEEE Trans. Neural Netw.* **2009**, *20*, 1368–1371. [\[CrossRef\]](#) [\[PubMed\]](#)
30. Mazzaglia, P.; Çatal, O.; Verbelen, T.; Dhoedt, B. Curiosity-Driven Exploration via Latent Bayesian Surprise. *arXiv* **2022**, arXiv:2104.07495. [\[CrossRef\]](#)
31. Schillaci, G.; Villalpando, A.P.; Hafner, V.V.; Hanappe, P.; Colliaux, D.; Wintz, T. Intrinsic motivation and episodic memories for robot exploration of high-dimensional sensory spaces. *Adapt. Behav.* **2020**, *29*, 549–566. [\[CrossRef\]](#)
32. Forestier, S.; Portelas, R.; Mollard, Y.; Oudeyer, P.Y. Intrinsically Motivated Goal Exploration Processes with Automatic Curriculum Learning. *J. Mach. Learn. Res.* **2022**, *23*, 1–41.
33. Storrs, K.R.; Kriegeskorte, N. Deep Learning for Cognitive Neuroscience. *arXiv* **2019**, arXiv:1903.01458. [\[CrossRef\]](#)
34. Macpherson, T.; Churchland, A.; Sejnowski, T.; DiCarlo, J.; Kamitani, Y.; Takahashi, H.; Hikida, T. Natural and Artificial Intelligence: A brief introduction to the interplay between AI and neuroscience research. *Neural Netw.* **2021**, *144*, 603–613. [\[CrossRef\]](#)
35. Piloto, L.S.; Weinstein, A.; Battaglia, P.; Botvinick, M. Intuitive physics learning in a deep-learning model inspired by developmental psychology. *Nat. Hum. Behav.* **2022**, *6*, 1257–1267. [\[CrossRef\]](#)
36. Huang, C.; Zhang, Z.; Mao, B.; Yao, X. An Overview of Artificial Intelligence Ethics. *IEEE Trans. Artif. Intell.* **2023**, *4*, 799–819. [\[CrossRef\]](#)
37. Utz, S.; Muscanell, N.L. Your Co-author Received 150 Citations: Pride, but Not Envy, Mediates the Effect of System-Generated Achievement Messages on Motivation. *Front. Psychol.* **2018**, *9*, 628. [\[CrossRef\]](#)
38. Gendolla, G.H.E. Surprise in the Context of Achievement: The Role of Outcome Valence and Importance. *Motiv. Emot.* **1997**, *21*, 165–193. [\[CrossRef\]](#)
39. Kiuru, N.; Spinath, B.; Clem, A.L.; Eklund, K.; Ahonen, T.; Hirvonen, R. The dynamics of motivation, emotion, and task performance in simulated achievement situations. *Learn. Individ. Differ.* **2020**, *80*, 101873. [\[CrossRef\]](#)
40. Deng, L. The mnist database of handwritten digit images for machine learning research. *IEEE Signal Process. Mag.* **2012**, *29*, 141–142. [\[CrossRef\]](#)
41. Simonyan, K.; Zisserman, A. Very Deep Convolutional Networks for Large-Scale Image Recognition. *arXiv* **2015**, arXiv:1409.1556. [\[CrossRef\]](#)
42. Kingma, D.P.; Ba, J. Adam: A Method for Stochastic Optimization. *arXiv* **2015**, arXiv:1412.6980.
43. O’Doherty, J.; Dayan, P.; Schultz, J.; Deichmann, R.; Friston, K.; Dolan, R.J. Dissociable Roles of Ventral and Dorsal Striatum in Instrumental Conditioning. *Science* **2004**, *304*, 452–454. [\[CrossRef\]](#)
44. Lillicrap, T.P.; Hunt, J.J.; Pritzel, A.; Heess, N.; Erez, T.; Tassa, Y.; Silver, D.; Wierstra, D. Continuous control with deep reinforcement learning. In Proceedings of the 4th International Conference on Learning Representations, San Juan, Puerto Rico, 2–4 May 2016; Bengio, Y., LeCun, Y., Eds.; ICLR: Vienna, Austria, 2016.
45. Sirigu, A.; Duhamel, J.R. Reward and decision processes in the brains of humans and nonhuman primates. *Dialogues Clin. Neurosci.* **2016**, *18*, 45–53. [\[CrossRef\]](#)
46. Morville, T.; Friston, K.; Burdakov, D.; Siebner, H.R.; Hulme, O.J. The Homeostatic Logic of Reward. *bioRxiv* **2018**. [\[CrossRef\]](#)
47. Hartwig, M.; Peters, A. Cooperation and Social Rules Emerging From the Principle of Surprise Minimization. *Front. Psychol.* **2021**, *11*, 606174. [\[CrossRef\]](#)
48. Montag, C.; Hahn, E.; Reuter, M.; Spinath, F.M.; Davis, K.; Panksepp, J. The Role of Nature and Nurture for Individual Differences in Primary Emotional Systems: Evidence from a Twin Study. *PLoS ONE* **2016**, *11*, e0151405. [\[CrossRef\]](#)

49. Zar, J.H. Spearman Rank Correlation: Overview, 2014. Wiley Online Library. Available online: <https://onlinelibrary.wiley.com/doi/10.1002/9781118445112.stat05964> (accessed on 15 November 2022).
50. Bi, X.Y.; Ma, X.; Abulaiti, A.; Yang, J.; Tao, Y. The influence of pride emotion on executive function: Evidence from ERP. *Brain Behav.* **2022**, *12*, e2678. [[CrossRef](#)] [[PubMed](#)]
51. Blanco, N.J.; Love, B.C.; Ramscar, M.; Otto, A.R.; Smayda, K.; Maddox, W.T. Exploratory decision-making as a function of lifelong experience, not cognitive decline. *J. Exp. Psychol. Gen.* **2016**, *145*, 284–297. [[CrossRef](#)] [[PubMed](#)]
52. Wykowska, A.; Chaminade, T.; Cheng, G. Embodied artificial agents for understanding human social cognition. *Phil. Trans. R. Soc. B* **2016**, *371*, 20150375. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.