

Statistical learning shapes pain perception and prediction independently of external cues

Reviewed Preprint

v2 • May 28, 2024

Revised by authors

Reviewed Preprint

v1 • October 4, 2023

Jakub Onysk , Nicholas Gregory, Mia Whitefield, Maeghal Jain, Georgia Turner, Ben Seymour, Flavia Mancini 

Computational and Biological Learning Unit, Department of Engineering, University of Cambridge, Cambridge CB2 1PZ, UK • Applied Computational Psychiatry Lab, Max Planck Centre for Computational Psychiatry and Ageing Research, Queen Square Institute of Neurology and Mental Health Neuroscience Department, Division of Psychiatry, University College London, UK • MRC Cognition and Brain Sciences Unit, University of Cambridge, UK • Wellcome Centre for Integrative Neuroimaging, John Radcliffe Hospital, Headington, Oxford OX3 9DU, UK • Center for Information and Neural Networks (CiNet), Osaka 565-0871, Japan

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

Abstract

The placebo and nocebo effects highlight the importance of expectations in modulating pain perception, but in everyday life we don't need an external source of information to form expectations about pain. The brain can learn to predict pain in a more fundamental way, simply by experiencing fluctuating, non-random streams of noxious inputs, and extracting their temporal regularities. This process is called statistical learning. Here we address a key open question: does statistical learning modulate pain perception? We asked 27 participants to both rate and predict pain intensity levels in sequences of fluctuating heat pain. Using a computational approach, we show that probabilistic expectations and confidence were used to weight pain perception and prediction. As such, this study goes beyond well-established conditioning paradigms associating non-pain cues with pain outcomes, and shows that statistical learning itself shapes pain experience. This finding opens a new path of research into the brain mechanisms of pain regulation, with relevance to chronic pain where it may be dysfunctional.

eLife assessment

This study presents a **valuable** insight into a computational mechanism of pain perception. The evidence supporting the authors' claims is **compelling**. The work will be of interest to pain researchers working on computational models and cognitive mechanisms of pain in a Bayesian framework.

<https://doi.org/10.7554/eLife.90634.2.sa2>

1 Introduction

Clinical pain typically varies over time; in most pain states, the brain receives a stream of volatile and noisy noxious signals, which are also auto-correlated in time. The temporal structure of these signals is important, because the human brain has evolved the exceptional ability to extract regularities from streams of auto-correlated sensory signals, a process called statistical learning [1–7]. In the context of pain, statistical learning can allow the brain to predict future pain, which is crucial for orienting behaviour and maximising well-being [8, 9]. Statistical learning might also be fundamental to the ability of the nervous system to endogenously regulate pain. Indeed, statistical learning generates predictions about forthcoming pain. We already know that pain expectations can modulate pain levels by gating the reciprocal transmission of neural signals between the brain and spinal cord, as shown by previous work on placebo and nocebo effects [10–14].

By using temporal sequences of noxious inputs, we have previously shown that the pain system supports the statistical learning of the basic rate of getting pain by engaging both somatosensory and supramodal cortical regions [8]. Specifically, both sensorimotor cortical regions and the ventral striatum encode probabilistic predictions about pain intensity, which are updated as a function of learning by engaging parietal and prefrontal regions. According to a Bayesian inference framework, both the predictive inference and its confidence should, in principle, modulate the neural response to noxious inputs and affect perception, as a function of learning. In support of this conjecture, there is evidence that the confidence of probabilistic pain predictions modulates the cortical response to pain [9]. The relationship is inverse: the lower the confidence, the higher is the early cortical response to noxious inputs (and viceversa), as measured by EEG. This is expected based on Bayesian inference theory: when confidence is low, the brain relies less on his prior beliefs and more on sensory evidence to respond to the input. Bayesian inference theory also predicts that prior expectations and their confidence scale perception [15]. Thus, we hypothesise that the predictions generated by learning the statistics of noxious inputs in dynamically evolving sequences of stimuli modulate the perception of forthcoming inputs.

Previously, it was found that pain perception is strongly influenced by probabilistic expectations as defined by a cue that predicts high or low pain [16]. In contrast to such cue-paradigm, the primary aim of our experiment was to determine whether the expectations participants hold about the sequence itself inform their perceptual beliefs about the intensity of the stimuli. To that end, we recruited 27 healthy participants to complete a psycho-physical experiment where we delivered four different, 80-trial-long sequences of evolving thermal stimuli, with four levels of temporal regularity. On each trial, a 2s thermal stimulus was applied, following which participants were asked to either rate their perception of the intensity (**Figure 1A**) or to predict the intensity of the next stimulus in the sequence (**Figure 1B**). Participants also reported their response confidence.

We contrasted four models of statistical learning, which varied according to the inference strategy used (i.e., optimal Bayesian inference or a heuristic) and the role of expectations on perception. All models used confidence ratings to weight the inference. We anticipate that probabilistic learning weighted by confidence and expectations modulates pain perception. This provides behavioural evidence for a link between learning and endogenous pain regulation. One reason why this is important is that it might help understand individual differences in the ability to endogenously regulate pain. This is particularly relevant for chronic pain, given that endogenous pain regulation can be dysfunctional in several chronic pain conditions [17–21], even before chronic pain

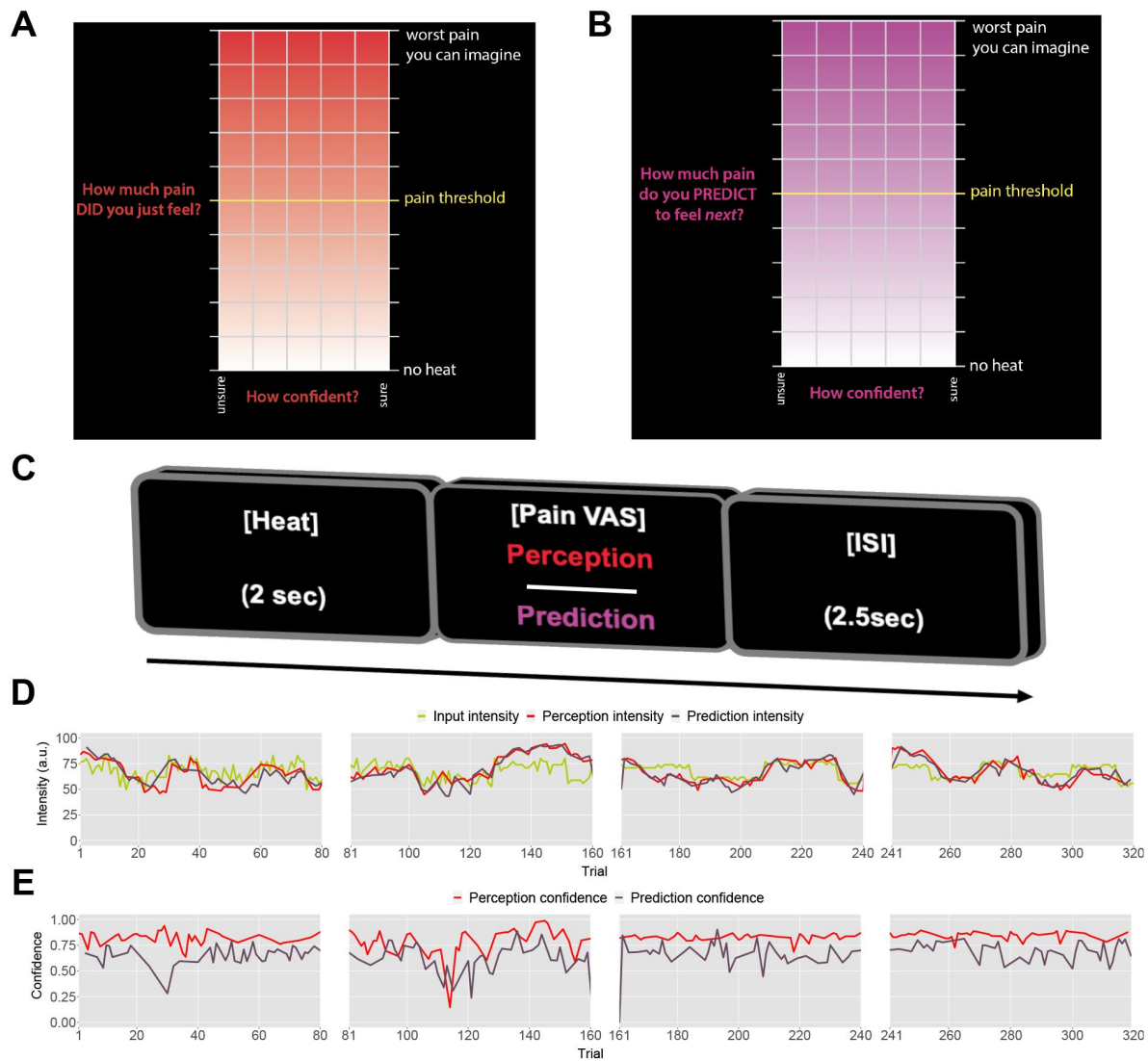


Figure 1.

Task design.

On each trial, each participant received a thermal stimulus lasting 2s from a sequence of intensities. This was followed by a perception (**A**) or a prediction (**B**) input screen, where the y-axis indicates the level of perceived/predicted intensity (0-100) centred around participant's pain threshold, and the x-axis indicates the level of confidence in one's perception (0-1). The inter-stimulus interval (ISI; black screen) lasted 2.5s (trial example in **C**). **D**: Example intensity sequences are plotted in green, participant's perception and prediction responses are in red and black. **E**: Participant's confidence rating for perception (red) and prediction (black) trials.

develops [22]. Although there is ample evidence for changes in the functional anatomy and connectivity of endogenous pain modulatory systems in chronic pain, their computational mechanisms are poorly understood.

2 Results

2.1 Model-naïve performance

Prior to modelling, we first checked whether participant's performance in the task was affected by the level of temporal regularity, i.e. the sequence condition. We varied the level of volatility and stochasticity across blocks (i.e. conditions), whilst we fixed their overall level within each block; the level of volatility was defined by the number of trials until the mean intensity level changes. The stochasticity is the additional noise that is added on each trial to the underlying mean, often referred to as the observation noise. The changes were often subtle and participants were not informed when they happened. We set two levels (low/high) of each type of uncertainty, achieving a 2x2 factorial design, with the order of conditions randomised across participants. A set of four example sequences of thermal intensities delivered to one of the participant's can be found in **Figure 1C**, alongside their ratings of perception and predictions. Additionally, example confidence ratings for each type of response are plotted in **Figure 1D**. Figures S2-3 show the plots of each participant's responses superimposed onto the sequences of noxious inputs.

As a measure of performance, we calculated the root mean squared error (RMSE) of participants responses (ratings and predictions) compared to the normative noxious input for each condition as in **Figure 2** (see also Methods). The lower the RMSE, the more accurate participants' responses are. Performance in different conditions was analysed with a repeated measures ANOVA, whose results are reported in full in Table S1. Although volatility did not affect rating accuracy ($F(1, 26) = 0.96, p = 0.336, \eta_p^2 = 0.036$), we found a 2-way interaction between the level of stochasticity of the sequence (low, high) and the type of rating provided (perceived intensity vs. prediction) ($F(1, 26) = 29.842, p < 0.001, \eta_p^2 = 0.534$). We followed up this interaction in post-hoc comparisons, as reported in Table S2. The performance score differences between all the pairs of stochasticity and response type interactions were significant, apart from the perception ratings in the stochastic environment as compared with perception and prediction performance in the low stochastic setting. Intuitively, the RMSE score analysis revealed an overall trend of participants performing worse on the prediction task, in particular when the

2.2 Modelling strategy

Our models were selected a-priori, following the modelling strategy from [23] and hence considered the same set of core models for clear extension of the analysis to our non-cue paradigm. The key question for us was whether expectations were used to weight the behavioural estimates during sequence learning. Therefore, we compared Bayesian and non-Bayesian models of sequential learning that weighted their ratings based on prior expectations versus two corresponding models that assumed perfect perception (i.e. not weighted by prior beliefs). As a baseline, we included a random-response model (please see Methods for a formal treatment of the computational models).

According to an optimal Bayesian inference strategy, on each trial, participants update their beliefs about the feature of interest (thermal stimuli) based on probabilistic inference, maintaining a full posterior distribution over its values [16, 24]. Operating within a Bayesian paradigm, participants are assumed to track and, following new information, update both the mean of the sequence of interest and the uncertainty around it [25]. In most cases, such inference makes an assumption about environmental dynamics. For example, a common assumption is that the

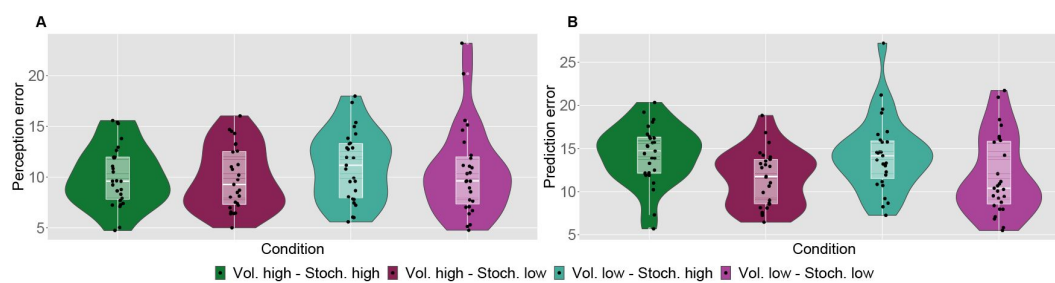


Figure 2.

Participant's model-naïve performance in the task. Violin plots of participant Root Mean Square Error (RMSE) for each condition for **A**: rating and **B**: prediction responses as compared with the input.

underlying mean (a hidden/latent state) evolves linearly according to a Gaussian random walk, with the rate of this evolution defined by the variance of this Gaussian walk (volatility). The observed value is then drawn from another Gaussian with that mean, which has some observation noise (stochasticity). In this case, the observer can infer the latent states through the process of Bayesian filtering [24], using the Kalman Filter (KF) algorithm [26] (Figure 3B).

Sequence learning can also be captured by a heuristic to the Bayesian inference, i.e. a simple reinforcement learning (RL) rule. Here, participants maintain and update a point-estimate of the expected value of the sequence in an adaptive manner, within a non-stationary environment. RL explicitly involves correcting the tracked mean of the sequence proportionally to a trial-by-trial prediction error (PE) - a difference between the expected and actual value of the sequence [27] (Figure 3A). Importantly, RL agents do not assume any specific dynamics of the environment and hence are considered model-free.

Both models perform a form of error correction about the underlying sequence. The rate at which this occurs is captured by the learning rate $\alpha \in [0, 1]$ element. The higher the learning rate, the faster participants update their beliefs about the sequence after each observation. For the RL model, the learning rate α is a free parameter that is constant across the trials. On the other hand, the learning rate in the KF model α_t (known as the Kalman gain) is calculated on every trial. It depends on participants' trial-wise belief uncertainty as well as their overall estimation of the inherent noise in the environment (stochasticity, s). The belief uncertainty is updated after each observation and depends on participants' sense of volatility (v) and stochasticity (s) in the environment.

Crucially, we also used participants' trial-by-trial confidence ratings to measure to what extent confidence plays a role in learning. This is captured by the confidence scaling factor C , which defines the extent to which confidence affects response (un-)certainty. Intuitively, the higher the confidence scaling factor C , the less important role confidence plays in participant's response. With relatively low values of C , when the confidence is low, participants' responses are more noisy, i.e. less certain. We demonstrate this in Figure 4 by plotting hypothetical responses (A-F) and the effect on the noise scaling (G-L) as a function of C and confidence ratings.

To evaluate the effect of expectation on perceived intensity (on top of statistical learning modulating perception), we expanded the standard RL and KF models by adding a perceptual weighting element, $\gamma \in [0, 1]$ (similarly to [16]). Essentially, γ governs how much each participant relies on the normative input on each trial, and how much their expectation of the input influences their reported perception - i.e. they take a weighted average of the two. The higher the γ , the bigger the impact of the expectation on perception. Again, in the case of the Reinforcement Learning model (eRL - expectation weighted RL), γ is a free parameter that is constant across trials, while in the Kalman Filter model (eKF - expectation weighted KF), γ_t is calculated on every trial and depends on: (1) the participants' trial-wise belief uncertainty, (2) their overall estimation of the inherent noise in the environment (stochasticity, s) and (3) the participant's subjective uncertainty about the level of intensity, ϵ .

Thus, in total we tested 5 models: RL and KF (perception not-weighted by expectations), eRL and eKF (perception weighted by expectations), and a baseline random model. We then proceeded to fit these 5 computational models to participants' responses. For parameter estimation, we used hierarchical Bayesian methods, where we obtained group- and individual-level estimates for each model parameter (see Methods).

2.3 Modelling results

We fit each model for each condition sequence. Example trial-by-trial model prediction plots from one participant can be found in Figure S4. To establish which of the models fitted the data best, we ran model comparison analysis based on the difference in expected log point-wise predictive

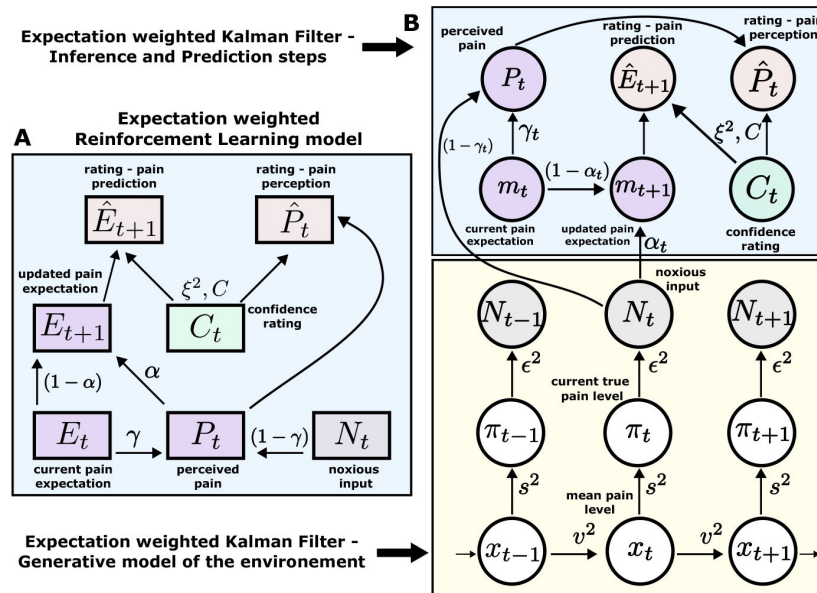


Figure 3.

Expectation weighted models.

Computational models used in the main analysis to capture participants' pain perception ($\hat{P}_t$) and prediction ($\hat{E}_{t+1}$) ratings. Both types of ratings are affected by confidence rating (C_t) on each trial. **(A)** In the Reinforcement Learning model, participant's pain perception (P_t) is taken to be weighted sum of the current noxious input (N_t) and their current pain expectation (E_t). Following the noxious input, participant updates their pain expectation (E_{t+1}). **(B)** In the Kalman Filter model, a generative model of the environment is assumed (yellow background) - where the mean pain level (x_t) evolves according to a Gaussian random walk (volatility v^2). The true pain level on each trial (π_t) is then drawn from a Gaussian (stochasticity s^2). Lastly, the noxious input, N_t , is assumed an imperfect indicator of the true pain level (subjective noise ϵ^2). Inference and prediction steps are depicted in a blue box. Participant's perceived pain is a weighted sum of expectation about the pain level (m_t) and current noxious input (N_t). Following each observation, N_t participant updates their expectation about the pain level (m_{t+1}).

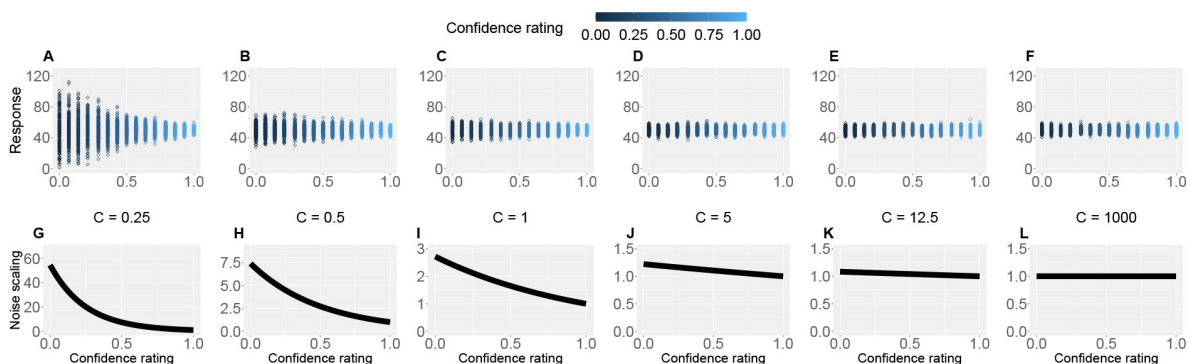


Figure 4.

Confidence scaling factor demonstration.

A-F: For a range of values of the confidence scaling factor C , we simulated a set of typical responses a participant would make for various levels of confidence ratings. The belief about the mean of the sequence is set at 50, while the response noise at 10. The confidence scaling factor C effectively scales the response noise, adding or reducing response uncertainty. **G-L:** The effect of different levels of parameter C on noise scaling. As C increases the effect of confidence is diminished.

density (ELPD) between models. The models are ranked according to the ELPD (with the largest providing the best fit). The ratio between the ELPD difference and the standard error around it provides a significance test proxy through the sigma effect. We considered at least a 2 sigma effect as indication of a significant difference. In each condition, the expectation weighted models (eKF and eRL) provided better fit than models without this element (KF and RL; approx. 2-4 sigma difference, as reported in **Figure 5A-D**) and Table S5. This suggests that regardless of the levels of volatility and stochasticity, participants still weigh perception of the stimuli with their expectation. In particular, we found that the expectation-weighted KF model offered a better fit than the eRL, although in conditions of high stochasticity this difference was short of significance against the eRL model. This suggests that in learning about temporal regularities in the sequences of thermal stimuli, participants' expectations modulate the perception of the stimulus. Moreover, this process was best captured by a model that updates the observer's belief about the mean and the uncertainty of the sequence in a Bayesian manner.

We also found that as the confidence in the response decreases, the response uncertainty is scaled linearly with a negative slope ranging between 0.112-0.276 across conditions (**Figure 6**), confirming the intuition that less confidence leads to bigger uncertainty.

As an additional check, for each participant, condition and response type (perception and prediction), we plotted participants' ratings against model predicted ratings and calculated a grand mean correlation in Figure S5. Next, we checked whether the parameters of the the winning eKF model differed across different sequence conditions. Given that volatility was fixed within condition, we treated it as a single-context scenario from the point of view of modelling [28], and we did not interpret its effect on the learning rate [29]. There were no differences for the group-level parameters; i.e., we did not detect significant differences between conditions in a hypothetical healthy participant group as generalised from our population of participants (Figure S12).

However, we found some differences at the individual-level of parameters (i.e. within our specific population of recruited participants), which we detected by performing repeated-measures ANOVAs (see Figure S13 for visualisation). The stochasticity parameter s was affected by the interaction between the levels of stochasticity and volatility ($F(1, 26) = 35.108, p < 0.001, \eta_p^2 = 0.575$), and was higher in highly stochastic and volatile conditions as compared to conditions where either volatility ($t = 7.735, p_{\text{bonf}} < 0.001$), stochasticity ($t = 9.396, p_{\text{bonf}} < 0.001$) or both were low ($t = 8.826, p_{\text{bonf}} < 0.001$). This suggests that, while participants' performance was generally worse in highly stochastic environments, participants seem to have attributed this to only one source - stochasticity (s), regardless of the source of higher uncertainty in the sequence (stochasticity or volatility).

The response noise ξ was modulated by the level of volatility ($F(1, 26) = 5.079, p = 0.033, \eta_p^2 = 0.163$), where it was smaller in highly volatile conditions. Moreover, we detected a significant interaction between volatility and stochasticity on the confidence scaling factor C ($F(1, 26) = 81.258, p < 0.001, \eta_p^2 = 0.758$), where the values C were overall lower when either volatility ($t = -11.570, p_{\text{bonf}} < 0.001$), stochasticity ($t = -6.165, p_{\text{bonf}} < 0.001$) or both ($t = -4.575, p_{\text{bonf}} < 0.001$) were high as compared to the conditions where both levels of noise were low. This indicates there may have been some trade-off between ξ and C , as lower values of C introduce additional uncertainty when participant's confidence is low.

Lastly, we found the initial uncertainty belief w_0 was affected by the interaction between volatility and stochasticity ($F(1, 26) = 5275.367, p < 0.001, \eta_p^2 = 0.995$) without a consistent pattern. All the other effects were not significant.

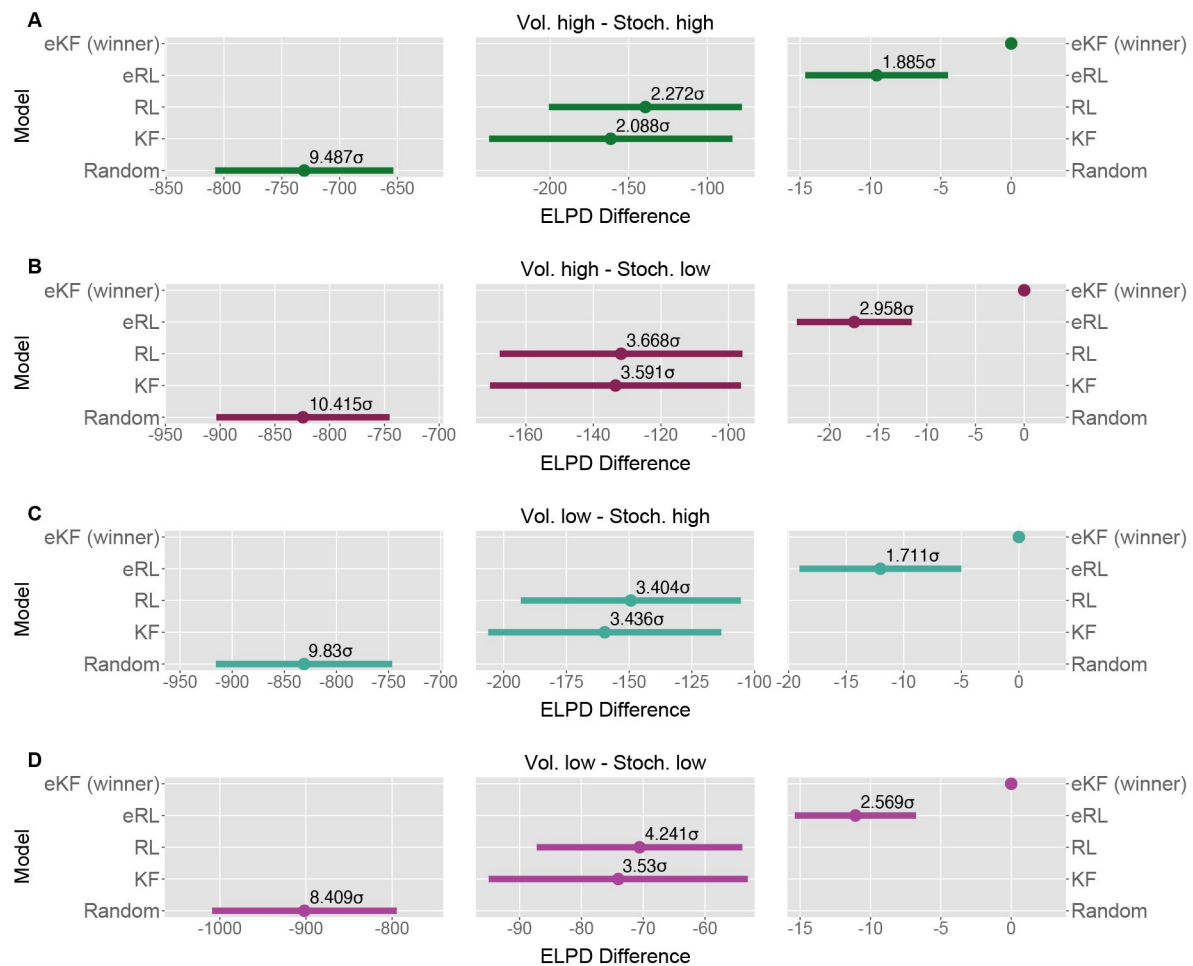


Figure 5.

Model comparison for each sequence condition (**A-D**). The dots indicate the ELPD difference between the winning model (eKF) every other model. The line indicates the standard error (SE) of the difference. The non-winning models' ELPD differences are annotated with the ratio between the ELPD difference and SE indicating the sigma effect, a significance heuristic.

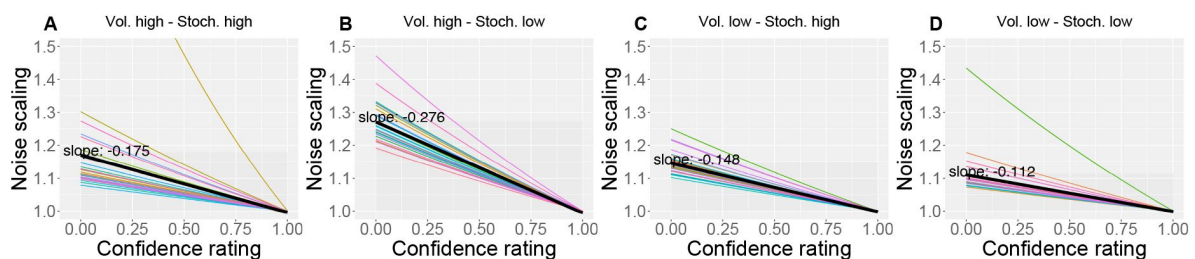


Figure 6.

(**A-D**): The effect of the confidence scaling factor on noise scaling for each condition. Each coloured line corresponds to one participant, with the black line indicating the mean across all participants. The mean slope for each condition is annotated.

In summary, we formalised the process behind pain perception and prediction in noxious time-series within the framework of sequential learning, where the best description of participants' statistical learning was captured through Bayesian filtering, in particular using a confidence-weighted Kalman Filter model. Most importantly, we discovered that, in addition to weighing their responses with confidence, participants used their expectations about stimulus intensity levels to form a judgement as to what they perceived. This mechanism was present across various levels of uncertainty that defined the sequences (volatility and stochasticity).

3 Discussion

Statistical learning allows the brain to extract regularities from streams of sensory inputs and is central to perception and cognitive function. Despite its fundamental role, it has often been overlooked in the field of pain research. Yet, chronic pain appears to fluctuate over time. For instance, [30–32] reported that chronic back pain ratings vary periodically, over several seconds-minutes and in absence of movements. This temporal aspect of pain is important because periodic temporal structures are easy to learn for the brain [1, 8]. If the temporal evolution of pain is learned, it can be used by the brain to regulate its responses to forthcoming pain, effectively shaping how much pain it experiences. Indeed, we show that healthy participants extract temporal regularities from sequences of noxious stimuli and use this probabilistic knowledge to form confidence-weighted judgements and predictions about the level of pain intensity they experience in the sequence. We formalised our results within a Bayesian inference framework, where the belief about the level of pain intensity is updated on each trial according to the amount of uncertainty participants ascribe to the stimuli and the environment. Importantly, their perception and prediction of pain were influenced by the expected level of intensity that participants held about the sequence before responding. When varying different levels of inherent uncertainty in the sequences of stimuli (stochasticity and volatility), the expectation and confidence weighted models fitted the data better than models weighted for confidence but not for expectations (Figure 5A–D). The expectation-weighted bayesian (KF) model offered a better fit than the expectation-weighted, model-free RL model, although in conditions of high stochasticity this difference was short of significance. Overall, this suggests that participants' expectations play a significant role in the perception of sequences of noxious stimuli.

3.1 Statistical inference and learning in pain sequences

The first main contribution of our work is towards the understanding of the phenomenon of statistical learning in the context of pain. Statistical learning is an important function that the brain employs across the lifespan, with relevance to perception, cognition and learning [6]. The large majority of past research on statistical learning focused on visual and auditory perception [1, 3, 4, 33], with the nociceptive system receiving relatively little attention [34]. Recently, we showed that the human brain can learn to predict a sequence of two pain levels (low and high) in a manner consistent with optimal Bayesian inference, by engaging sensorimotor regions, parietal, premotor regions and dorsal striatum [8]. We also found that the confidence of these probabilistic inferences modulates the cortical response to pain, as expected by hierarchical Bayesian inference theory [9]. Here we tested sequences with a much larger range of stimulus intensities to elucidate the effect of statistical learning and expectations on pain perception. As predicted by hierarchical Bayesian inference theory, we find that the pain intensity judgements are scaled by both expectations and confidence.

Hence, our work highlights the inferential nature of the nociceptive system [16, 34–37], where in addition to the sheer input received by the nociceptors, there is a wealth of a priori knowledge and beliefs the agent holds about themselves and the environment that need to be

integrated to form a judgement about pain [38–41]. This has an immediate significance for the real world, where weights need to be assigned to prior beliefs and/or stimuli to successfully protect the organism from further damage, but only to an extent to which it is beneficial.

Secondly, our results regarding the effect of expectation on pain perception relate to a much larger literature on this topic. The prime example would be placebo analgesia (i.e. the expectation of pain relief decreasing pain perception) and nocebo hyperalgesia (i.e. the expectation of high level of pain increasing its perception; [10, 37, 42, 43]). Recent work attempted to capture such expectancy effects within the Bayesian inference framework. For example, [25] showed that in addition to expectation influencing perceived pain in general, higher level of uncertainty around that expectation attenuated its effect on perception. Similarly, [44] demonstrated that when the discrepancy between the expectation and outcome (prediction error) is unusually large, the role of expectation is significantly reduced and so the placebo and nocebo effects are not that strong. An unusually large prediction error could be thought of as contributing to increased uncertainty about the stimuli, which mirrors the results from [25] Bayesian framework. Nevertheless, the types of stimuli used in the above studies (i.e. noxious stimuli cued by non-noxious stimuli) differed from the more ecologically valid sequences of pain that are reported by chronic pain patients [30], as we indicated above. Furthermore, [16] used a conditioning paradigm and also found that expectations influence both perception and learning, in a self-reinforcing loop. Our work has followed a similar modelling strategy to [16], but it goes beyond simple conditioning schedules or sequences of two-level discrete painful stimuli, showing expectancy effects even when the intensities are allowed to vary across a wider range of values and according to more complex statistical temporal structures. Additionally, given the reported role of confidence in the perception of pain [9, 45], we draw a more complete picture by including participants confidence ratings in our modelling analysis.

Future studies would need to determine whether statistical learning and its effect on pain is altered in chronic pain conditions. This is important because statistical learning could, in principle, influence how a pain state evolves. Once a pain state is initiated, how an individual learns and anticipates the fluctuating pain signals may contribute to determine how well it can be regulated by the nervous system, thus affecting the severity and recurrence of pain flares. This, in turn, would affect whether aversive associations with the instigating stimulus are extinguished or reinforced [36]. In chronic pain, dysfunctional learning may promote the amplification and maintenance of pain signals, contributing to the reinforcement of aversive associations with incident stimuli, as well as the persistence of pain [46–48].

Our paper comes with open tools, which can be adapted in future studies on statistical learning in chronic pain. The key advantage of taking an hypothesis-driven, computational-neuroscience approach to quantify learning is that it allows to go beyond symptoms-mapping, identifying the quantifiable computational principles that mediate the link between symptoms and neural function.

In summary, we show that statistical expectations and confidence scale the judgement of pain in sequences of noxious stimuli as predicted by hierarchical Bayesian inference theory, opening a new avenue of research on the role of learning in pain.

4 Materials and methods

4.1 Participants

Thirty-three (18 female) healthy adult participants were recruited for the experiment. The mean age of participants was 22.4 ± 2.7 years old (range: 18–35). Participants had no chronic condition and no infectious illnesses, as well as no skin conditions (e.g. eczema) at the site of stimulus

delivery. Moreover, we only recruited participants that had not taken any anti-anxiety, anti-depressive medication, nor any illicit substances, alcohol and pain medication (including NSAIDs such as ibuprofen and paracetamol) in the 24 hours prior to the experiment. All participants gave informed written consent to take part in the study, which was approved by the local ethics committee.

4.2 Protocol

The experimental room's temperature was maintained between 20°C to 23°C. Upon entry, an infrared thermometer was used to ensure participants temperature was above 36°C at the forehead and forearm of the non-dominant hand, to account for the known effects of temperature on pain perception [49]. A series of slideshows were presented, which explained the premise of the experiment and demonstrated what the participant would be asked to carry out. Throughout this presentation, questions were asked to ensure participants understood the task. Participants were given multiple opportunities to ask questions throughout the presentation.

We used the Medoc Advanced Thermosensory Stimulator 2 (TSA2) [50] to deliver thermal stimuli using the CHEPS thermode. The CHEPS thermode allowed for rapid cooling (40°C / sec) and heating (70°C / sec) so transitions between the baseline and stimuli temperatures were minimal. The TSA2 was controlled externally, via Matlab (Mathworks).

We then established the pain threshold, using the method of limits [51], in order to centre the range of temperature intensities used in the experiment. Each participant was provided with stimuli of increasing temperature, starting from 40°C going up in 0.5°C increments, using an inter-stimulus interval (ISI) of 2.5 sec and a 2 sec duration. The participant was asked to indicate when the stimuli went from warm to painful - this temperature was noted and the stimuli ended. The procedure was repeated three times, and the average was used as an estimate of the pain threshold.

During the experiment, four sequences of thermal stimuli were delivered. Due to the phenomenon of offset analgesia (OA), where decreases in tonic pain result in a proportionally larger decrease in perceived pain [52], we chose phasic stimuli, with a duration of 2 sec and an inter-stimulus-interval (ISI) 2.5 seconds. In order to account for individual differences, the temperatures which the levels refer to are based upon the participants threshold.

The median intensity level was defined as threshold, giving a max temperature of 3°C above threshold, which was found to be acceptable by participants. Before the start of the experiment each participant was provided with the highest temperature stimuli that could be presented, given their measured threshold, to ensure they were comfortable with this. Two participants found the stimulus too painful - the temperature range was lowered by 1°C and this was found to be acceptable.

After every trial of each sequence, the participant was asked for either their perception of the previous stimulus, or their prediction for the next stimulus through a 2D VAS (**Figure 1B**), presented using PsychToolBox-3 [53]. The y-axis encodes the intensity of the stimulus either perceived or predicted, ranging from 0 (no heat detected/predicted) to 100 (worst pain imaginable perceived/predicted); on this scale, 50 represents the pain threshold. This was done as a given sequence was centred around the threshold. The x-axis encodes confidence in either perception or prediction, ranging from 0 - completely uncertain ('unsure') - to 1 - complete confidence in the rating ('sure'). Differing background colours were chosen to ensure participants were aware of what was being asked, and throughout the experiment participants were reminded to take care in answering the right question. The mouse movement was limited to be inside of the coloured box, which defined the area of participants' input. At the beginning of each input screen, the mouse location was uniformly randomised within the input box.

The sequence of response types was randomised so as to retain 40 prediction and 40 perception ratings for each of the four sequence conditions. For an 80-trial long sequence, this gave 80 participant responses. Each sequence condition was separated by a 5 minute break, during which the thermode's probe was slightly moved around the area of skin on the forearm to reduce sensitisation (i.e. a gradual increase in perceived intensity with repetitive noxious stimuli) [54]. In the middle of each sequence, there was a 3 minute break. During the ISI, the temperature returned to a base-line of 38°C. One participant was unable to complete the sequence as their threshold was too low, and data from four participants was lost due to Medoc software issues (the remote control failed and the data of 2 out of 4 sessions were not saved). We excluded one participant's whose ratings/predictions were inversely proportional to the noxious input. Thus, we analysed data from 27 participants.

4.2.1 Generative process of the painful sequences

We manipulated two sources of uncertainty in the sequence: the stochasticity (s) of the observation and the volatility (v) of the underlying sequence [29]. Sequences were defined by two levels (high or low) of stochasticity and volatility, resulting in four different sequences conditions - creating a 2x2 factorial design. Each sequence was defined as a series of chunks, where the intensity for trial t , i_t was sampled from $\mathcal{N}(I, \sigma^2)$, where σ^2 indicates the level of stochasticity ($\sigma^2 = 1.75$ for high level of stochasticity, $\sigma^2 = 0.25$ for low level of stochasticity). The mean of the chunk, I , was drawn from $\mathcal{N}(3.5, 10.5)$. To ensure a noticeable difference in chunk intensity to the participant, concurrent chunk means were constrained to be at least 2 intensity levels different. Volatility was implemented by defining the length, or number of trials, of a chunk (l) drawn from $\mathcal{N}(L, a, L + a)$, where L is the mean of the chunk length ($L = 15$ for high volatility level, $L = 25$ for low volatility level). A jitter, a , was added around the mean to ensure the transition from one chunk to the next was not consistent or predictable. For both high and low volatility conditions we set $a = 3$. Sampled values were then discretised, where any intensities outside the valid intensity range [1, 13] were discarded and re-sampled resulting in an 80-trial long sequence for each condition. The mean of each sequence was centred around intensity level 7, i.e. the participants threshold. So defined, six sets of four sequences were sampled. Each participant received one set, with a randomised sequence order. See an example sequence (after subject-specific linear transformation) and one participant's responses (including confidence ratings) in Figure 1C-D.

4.3 Data pre-processing

Since the intensity values of the noxious input were discretised between 1 and 13, while the participant's responses (perception and prediction) were given on a 0-100 scale, we applied a linear transformation of the input to map its values onto a common 0-100 range. For each participant, for a set of inputs at perception trials from the concatenated sequence (separate sequence conditions in the order as presented), we fit a linear least-squares regression using Python's `scipy.stats.linregress` function. On rare occasions, when the transformed input was negative, we refit the line using Python's non-linear least squares function `scipy.optimize.curve_fit`, constraining the intercept above 0 [55]. We then extracted each participant's optimised slope and intercept and applied the transformation both to the concatenated and condition-specific sequence of inputs. So transformed, the sequences were then used in all the analyses. Plots of each participant transformation can be found in Figure S1. We superimposed participant's responses onto the noxious input condition sequences in Figure S2.

To capture participant's model-naive performance in the task, both for the concatenated and condition-specific sequence, we calculated Root Mean Square Error (RMSE) of each participant's perception (Eq. 1) and pre-diction (Eq. 2) responses as compared to the input. The lower the RMSE, the higher the response accuracy.

$$RMSE_P = \sqrt{\frac{\sum_{t=1}^{T_P} (y_t - \hat{P}_t)^2}{T_P}} \quad (\text{Eq. 1})$$

$$RMSE_E = \sqrt{\frac{\sum_{t=1}^{T_E} (y_{t+1} - \hat{E}_{t+1})^2}{T_E}} \quad (\text{Eq. 2})$$

where T_P is the number of perception trials, $\hat{P}_t$ is participant's perception response to the stimulus y_t at trial t , T_E is the number of prediction trials and $\hat{E}_{t+1}$ is participant's prediction of the next stimulus intensity y_{t+1} at trial $t + 1$.

4.4 Models

4.4.1 Reinforcement Learning

RI

In reinforcement learning models, learning is driven by discrepancies between the estimate of the expected value and observed values. Before any learning begins, at trial $t = 1$, participants have an initial expectation, $E_1 = E_0$, which is a free parameter that we estimate.

On each trial, participants receive a thermal input N_t . We then calculate the prediction error δ_t , defined as the difference between the expectation E_t and the input N_t (Eq. 3).

$$\delta_t = N_t - E_t \quad (\text{Eq. 3})$$

Participant is then assumed to update their expectation of the stimulus on the next trial as in Eq. 4

$$E_{t+1} = E_t + \alpha \delta_t \quad (\text{Eq. 4})$$

where α is the learning rate (free parameter), which governs how fast participants assimilate new information to update their belief.

On trials when participants rate their perceived intensity, we assume no effects on their perception other than confidence rating c_t and response noise, so participants perception response $\hat{P}_t$ is drawn from a Gaussian distribution, with the mean $P_t = N_t$ and a confidence-scaled response noise ξ (free parameter), as in Eq. 5

$$\hat{P}_t \sim \mathcal{N}\left(P_t, \xi^2 \exp\{C^{-1}(1 - c_t)\}^2\right) \quad (\text{Eq. 5})$$

where C is the confidence scaling factor (free parameter), which defines the extent to which confidence affects response uncertainty. Please, see **Figure 4** for an intuition behind confidence scaling.

On trials when participants are asked to predict the intensity of the next thermal stimulus, we use the updated expectation E_{t+1} to model participants prediction response $\hat{E}_{t+1}$. This is similarly affected by confidence rating and response noise and is defined as in [Eq. 6](#).

$$\hat{E}_{t+1} \sim \mathcal{N}\left(E_{t+1}, \xi^2 \exp\{C^{-1}(1 - c_t)\}^2\right) \quad (\text{Eq. 6})$$

To recap, the RL model has 4 free parameters: the learning rate α , response noise ξ , the initial expectation E_0 , and the confidence scaling factor C .

eRL

Additionally, where we investigate the effects of expectation on the perception of pain [[16](#)], we included an element that allows us express the perception as a weighted sum of the input and expectation ([Eq. 7](#)).

$$P_t = (1 - \gamma)N_t + \gamma E_t \quad (\text{Eq. 7})$$

where γ [[0, 1](#)] (free parameter) captures how much participants rely on the normative thermal input vs. their expectation. When $\gamma = 0$, the expectation plays no role and the model simplifies to that of the standard RL above. In total, the eRL model has 5 free parameters, with the other equations the same as in the RL model, with the exception of the prediction error, which now relies on the expectation-weighted pain perception P_t ([Eq. 8](#)).

$$\delta_t = P_t - E_t \quad (\text{Eq. 8})$$

4.4.2 Kalman Filter

Kf

To capture sequential learning in a Bayesian manner, we used the Kalman filter model [[16](#), [24](#), [26](#)]. KF assumes a generative model of the environment where the latent state on trial t , x_t (the mean of the sequences in the experiment), evolves according to a Gaussian random walk with a fixed drift rate, v (volatility), as in [Eq. 9](#).

$$x_t \sim \mathcal{N}(x_{t-1}, v^2) \quad (\text{Eq. 9})$$

The observation on trial t , N_t , is then drawn from a Gaussian ([Eq. 10](#)) with a fixed variance, which represents the observation uncertainty s (stochasticity).

$$N_t \sim \mathcal{N}(x_t, s^2) \quad (\text{Eq. 10})$$

As such the KF assumes stable dynamics since the generative process has fixed volatility and stochasticity.

For ease of explanation, we refer to the thermal input at each trial as N_t , we also use the $N_{1:t}$ notation, which refers to a sequence of observations up to and including trial t . The model allows to obtain posterior beliefs about the latent state x_t given the observations. This is done by tracking an internal estimate of the mean m_t and the uncertainty, w_t , of the latent state x_t .

First, following standard KF results, on each trial, the participant is assumed to hold a prior belief (indicated with $(-)$ superscript) about the latent state, x_t ([Eq. 11](#)).

$$x_t | N_{1:t-1} \sim \mathcal{N}\left(m_t^{(-)}, w_t^{2(-)}\right) \quad (\text{Eq. 11})$$

On the first trial, before any observations, we set $m_1^{(-)} = E^0, w_1^{(-)} = w_0$ (free parameters). In light of the new observation, N_t on trial t , the tracked mean and uncertainty of the latent state are reweighed based on the new evidence N_t and its associated observation uncertainty s as in Eq. 12.

$$x_t | N_{1:t} \sim \mathcal{N} \left(\frac{s^2 m_t^{(-)} + w_t^{2(-)} N_t}{s^2 + w_t^{2(-)}}, \frac{s^2 w_t^{2(-)}}{s^2 + w_t^{2(-)}} \right) \quad (\text{Eq. 12})$$

We can then define the learning rate α_t (Eq. 13),

$$\alpha_t = \frac{w_t^{2(-)}}{s^2 + w_t^{2(-)}} \quad (\text{Eq. 13})$$

to get the update rule for the new posterior beliefs (indicated with $(+)$ superscript) about the mean (Eq. 14) and uncertainty (Eq. 15) of x_t .

$$m_t^{(+)} = m_t^{(-)}(1 - \alpha_t) + N_t \alpha_t \quad (\text{Eq. 14})$$

$$w_t^{2(+)} = w_t^{2(-)}(1 - \alpha_t) \quad (\text{Eq. 15})$$

Following this new belief, and the assumption about the environmental dynamics (volatility), the participant forms a new prior belief about the latent state x_{t+1} for the next trial $t + 1$ as in Eq. 16.

$$x_{t+1} | N_{1:t} \sim \mathcal{N} \left(m_{t+1}^{(-)}, w_{t+1}^{2(-)} \right) \quad (\text{Eq. 16})$$

Where

$$m_{t+1}^{(-)} = m_t^{(+)} \quad (\text{Eq. 17})$$

$$w_{t+1}^{2(-)} = w_t^{2(+)} + v^2 \quad (\text{Eq. 18})$$

We can simplify the notation to make it comparable to the RL models. We let, $m_{t+1}^{(+)} = m_{t+1}^{(-)}$, and $w_{t+1}^{2(-)} = w_{t+1}^{2(+)} + v^2$. Following a new observation at trial t , we calculate the prediction error (Eq. 19) and learning rate (Eq. 20).

$$\delta_t = y_t - m_t \quad (\text{Eq. 19})$$

$$\alpha_t = \frac{w_t^2}{w_t^2 + s^2} \quad (\text{Eq. 20})$$

we then update the belief about the mean (Eq. 21) and uncertainty (Eq. 22) of the latent state for the next trial.

$$\begin{aligned} m_{t+1} &= m_t(1 - \alpha_t) + N_t \alpha_t \\ &= m_t + \alpha_t(N_t - m_t) \end{aligned} \quad (\text{Eq. 21})$$

$$w_{t+1}^2 = w_t^2(1 - \alpha_t) + v^2 \quad (\text{Eq. 22})$$

Now, mapping this onto the experiment, the mean of the latent state is participants expectation $E_t = m_t$, and so we have participant perception rating modelled as in [Eq. 23](#).

$$\hat{P}_t \sim \mathcal{N}\left(P_t, \xi^2 \exp\{C^{-1}(1 - c_t)\}^2\right) \quad (\text{Eq. 23})$$

and the prediction rating for the next trial as in [Eq. 24](#).

$$\hat{E}_{t+1} \sim \mathcal{N}\left(E_{t+1}, \xi^2 \exp\{C^{-1}(1 - c_t)\}^2\right) \quad (\text{Eq. 24})$$

In total the model has 6 free parameters: s (environmental stochasticity), v (environmental volatility), ξ (response noise), E_0 (initial belief about the mean), w_0 (initial belief about the uncertainty) and C (confidence scaling factor).

eKF - expectation weighted Kalman Filter

We can introduce the effect of expectation on the pain perception, by assuming that participants treat the thermal input as an imperfect indicator of the true level of pain [[16](#)]. In this case, the input, N_t , is modelled as in [Eq. 25](#)

$$N_t \sim \mathcal{N}(\pi_t, \varepsilon^2) \quad (\text{Eq. 25})$$

which forms an expression for the likelihood of the observation and adds an additional level to the inference, slightly modifying the Kalman filter assumptions such that:

$$\pi_t \sim \mathcal{N}(x_t, s^2) \quad (\text{Eq. 26})$$

However, we can apply the standard KF results and Bayes' rule to arrive at simple update rules for the participants' belief about the mean and uncertainty of the latent state x_t . From this, we get a prior on the π_t defined in [Eq. 27](#)

$$\pi_t | N_{1:t-1} \sim \mathcal{N}\left(m_t^{(-)}, w_t^{2(-)} + s^2\right) \quad (\text{Eq. 27})$$

which, following a new input N_t , gives us the posterior belief about π_t as in [Eq. 28](#).

$$\pi_t | N_{1:t} \sim \mathcal{N}\left(\frac{\varepsilon^2 m_t^{(-)} + (s^2 + w_t^{2(-)}) N_t}{\varepsilon^2 + s^2 + w_t^{2(-)}}, \frac{\varepsilon^2 (s^2 + w_t^{2(-)})}{\varepsilon^2 + s^2 + w_t^{2(-)}}\right) \quad (\text{Eq. 28})$$

Now, if we define γ_t as in [Eq. 29](#)

$$\gamma_t = \frac{\varepsilon^2}{\varepsilon^2 + s^2 + w_t^{2(-)}} \quad (\text{Eq. 29})$$

we have that the posterior belief about the mean level of pain π_t is calculated as:

$$P_t^{(+)} = \gamma_t m_t^{(-)} + (1 - \gamma_t) N_t \quad (\text{Eq. 30})$$

which is a weighted sum of the input N_t and participant expectation about the latent state x_t , governed by the perceptual weight γ_t , analogously to the eRL model. Finally, the posterior belief about x_t is obtained in Eq. 31.

$$x_t | N_{1:t} \sim \mathcal{N} \left(\frac{(\varepsilon^2 + s^2)m_t^{(-)} + w_t^{2(-)}N_t}{\varepsilon^2 + s^2 + w_t^{2(-)}}, \frac{(\varepsilon^2 + s^2)w_t^{2(-)}}{\varepsilon^2 + s^2 + w_t^{2(-)}} \right) \quad (\text{Eq. 31})$$

Now, setting the learning rate as in Eq. 32

$$\alpha_t = \frac{w_t^2}{\varepsilon^2 + w_t^2 + s^2} \quad (\text{Eq. 32})$$

we get:

$$m_t^{(+)} = m_t^{(-)}(1 - \alpha_t) + N_t \alpha_t \quad (\text{Eq. 33})$$

$$w_t^{2(+)} = w_t^{2(-)}(1 - \alpha_t) \quad (\text{Eq. 34})$$

Next, following the same notation simplification as before, we get the update rules for the prior belief about the mean (Eq. 35) and uncertainty (Eq. 36) of the latent state x_{t+1} for the next trial.

$$\begin{aligned} m_{t+1} &= m_t(1 - \alpha_t) + N_t \alpha_t \\ &= m_t + \alpha_t(N_t - m_t) \end{aligned} \quad (\text{Eq. 35})$$

$$w_{t+1}^2 = w_t^2(1 - \alpha_t) + v^2 \quad (\text{Eq. 36})$$

as well as the expression for subjective perception, P_t , at trial t (Eq. 37).

$$P_t = \gamma_t m_t + (1 - \gamma_t) N_t \quad (\text{Eq. 37})$$

The perception and prediction responses are modelled analogously as the KF model. In total, the model has 7 free parameters: ε (subjective noise), s (environmental stochasticity), v (environmental volatility), ξ (response noise), E_0 (initial belief about the mean), w_0 (initial belief about the uncertainty) and C (confidence scaling factor).

4.4.3 Random model

As a baseline, we also included a model that performs a random guess. The perceptual/prediction ratings were modelled as in Eq. 38.

$$\hat{P}_t \sim \mathcal{N} \left(R, \xi^2 \exp \{ C^{-1} (1 - c_t) \}^2 \right) \quad (\text{Eq. 38})$$

The model has 3 free parameters: R , ξ , and C , where R is a constant value that participants respond with.

4.5 Model fitting

Model parameters were estimated using hierarchical Bayesian methods, performed with RStan package (v. 2.21.0) [56] in R (v. 4.0.2) based on Markov Chain Monte Carlo techniques (No-U-Turn Hamiltonian Monte Carlo). For the individual level-parameters we used non-centred parametrisation [57]. For the group-level parameters we used $\mathcal{N}(0, 1)$ priors for the mean, and

the gamma-mixture representation of the Student-t(3,0,1) for the scale [58]. Parameters in the (0, 1) range were constrained using Phi_approx - a logistic approximation to the cumulative Normal distribution [59].

For each condition and each of the four chains, we ran 6000 samples (after discarding 6000 warm-up ones). For each condition, we examined R-hat values for each individual- (including the $\mathcal{N}(0, 1)$ error term from the non-centred parametrisation) and group-level parameters from each model to verify whether the Markov chains have converged. At the group-level and individual-level, all R-hat values had a value < 1.1 , indicating convergence. In the random response model, 0.01% – 0.16% iterations saturated the maximum tree depth of 11.

4.5.1 Model comparison

For model comparison, we used R package loo, which provides efficient approximate leave-one-out cross-validation. The package allows to estimate the difference in models' expected predictive accuracy through the difference in expected log point-wise predictive density (ELPD) [60]. By looking at the ratio between the ELPD difference and the standard error (SE) of the difference, we get the sigma effect - a heuristic for significance of such model differences. There's no agreed-upon threshold of SEs that determines significance, but the higher the sigma difference, the more robust is the effect. The closeness of fit can be also captured with LOO information criterion (LOOIC), where the lower LOOIC values indicate better fit.

4.5.2 Parameter comparison

For the comparison of group-level parameters between conditions, we extracted 95% High Density Intervals (HDI) of the permuted and merged (across chains) posterior samples of each group-level parameter [61]. To assess significant differences between conditions, we calculated a difference between such defined intervals. In the Bayesian scenario, a significant difference is indicated by the interval not containing the value 0 [62, 63].

4.5.3 Parameter and model recovery

To assess the reliability of our modelling analysis [64], for each model we performed parameter recovery analysis, where we simulated participants' responses using newly drawn individual-level parameters from the group-level distributions.

We repurposed existing sequences of noxious inputs in the [1, 13] range (pre-transformation). When then applied a linear transformation to the input sequences using sampled slope and intercept coefficients from a Gaussian distribution of these coefficients that we estimated based on our dataset using R's fitdistrplus package. Furthermore, we simulated the confidence ratings based on lag-1 auto-correlation across a moving window of the transformed input sequence.

We then fit the same model to the simulated data and calculated Pearson correlation coefficients r between the generated and estimated individual-level parameters. The higher the coefficient r , the more reliable the estimates are, which can be categorised as: poor (if $r < 0.5$); fair (if $0.5 < r < 0.75$); good ($0.75 < r < 0.9$); excellent (if $r > 0.9$) [65]. Results are reported in Table S3 and Figures S6-11.

We also performed model recovery analysis [64], where we first simulated responses using each model and then fit each model-specific dataset with each model. We then counted the number of times a model fit the simulated data best (according to the LOOIC rule), effectively creating an $M \times M$ confusion matrix, where M is the number of models. In the case where we have a diagonal matrix of ones, the models are perfectly recoverable and hence as reliable as possible. Results are reported in Table S4.

In Tables S6-S9 we report Bulk and Tail Effective Sample Size (ESS) for each condition, for each model and parameter.

Data Availability

All code and data will be available open source, released upon acceptance of the paper in a peer-reviewed journal.

5 Data and code availability

All code and data will be available open source, released upon acceptance of the paper.

Acknowledgements

The study was funded by a MRC Career Development Award to FM (MR/T010614/1) and a UKRI Advanced Pain Discovery Platform grant to both F.M. and B.S. (MR/W027593/1). B.S. was also funded by Wellcome (214251/Z/18/Z), Versus Arthritis (21537), and IITP (MSIT 2019-0-01371). This work has been performed using resources provided by the Cambridge Tier-2 system operated by the University of Cambridge Research Computing Service (www.hpc.cam.ac.uk) funded by EPSRC Tier-2 capital grant (EP/T022159/1). HPC access was additionally funded by an EPSRC research infrastructure grant to F.M.. We are grateful to Professor Máté Lengyel and Professor Deborah Talmi for helpful discussions about the study. For the purpose of open access, the author has applied a Creative Commons Attribution (CC BY) licence to any Author Accepted Manuscript version arising from this submission.

7 Author contributions

JO, NG, MW, GT, BS and FM designed the study. NG collected the data. JO, NG, MJ and FM analysed the data. JO, NG, MW, BS and FM wrote the article.

8 Declaration of interests

The authors declare no competing interests.

References

1. Dehaene S., Meyniel F., Wacongne C., Wang L., Pallier C. (2015) **The Neural Representation of Sequences: From Transition Probabilities to Algebraic Patterns and Linguistic Trees** *Neuron* **88**:2–19
2. Schapiro A., Turk-Browne N. (2015) **Statistical learning** *Brain mapping* **3**:501–506
3. Fiser J., Lengyel G. (2019) **A common probabilistic framework for perceptual and statistical learning** *Current Opinion in Neurobiology* **58**:218–228
4. Meyniel F., Maheu M., Dehaene S. (2016) **Human Inferences about Sequences: A Minimal Transition Probability Model** *PLOS Computational Biology* **12**
5. Kourtzi Z., Welchman A. E. (2019) **Learning predictive structure without a teacher: decision strategies and brain routes** *Current opinion in neurobiology* **58**:130–134
6. Sherman B. E., Graves K. N., Turk-Browne N. B. (2020) **The prevalence and importance of statistical learning in human cognition and behavior** *Current Opinion in Behavioral Sciences* **32**:15–20
7. Turk-Browne N. B., Scholl B. J., Chun M. M., Johnson M. K. (2009) **Neural Evidence of Statistical Learning: Efficient Detection of Visual Regularities Without Awareness** *Journal of Cognitive Neuroscience* :1934–1945
8. Mancini F., Zhang S., Seymour B. (2022) **Computational and neural mechanisms of statistical pain learning** *Nature Communications* **13**
9. Mulders D., Seymour B., Mouraux A., Mancini F. (2023) **Confidence of probabilistic predictions modulates the cortical response to pain** *Proceedings of the National Academy of Sciences* **120**
10. Tracey I. (2010) **Getting the pain you expect: mechanisms of placebo, nocebo and reappraisal effects in humans** *Nature Medicine* **16**:1277–1283
11. Tinnermann A., Geuter S., Sprenger C., Finsterbusch J., Büchel C. (2017) **Interactions between brain and spinal cord mediate value effects in nocebo hyperalgesia** *Science* :105–108
12. Eippert F., Finsterbusch J., Bingel U., Büchel C. (2009) **Direct evidence for spinal cord involvement in placebo analgesia** *Science* **326**:404–404
13. Geuter S., Büchel C. (2013) **Facilitation of pain in the human spinal cord by nocebo treatment** *Journal of Neuroscience* **33**:13784–13790
14. Fields H. L. (2018) **How expectations influence pain** *Pain* **159**:S3–S10
15. Knill D. C., Richards W. (1996) **Perception as Bayesian inference**
16. Jepma M., Koban L., van Doorn J., Jones M., Wager T. D. (2018) **Behavioural and neural evidence for self-reinforcing expectancy effects on pain** *Nature Human Behaviour* :838–855

17. Bushnell M. C., Čeko M., Low L. A. (2013) **Cognitive and emotional control of pain and its disruption in chronic pain** *Nature Reviews Neuroscience* **14**:502–511
18. Bruehl S., McCubbin J. A., Harden R. N. (1999) **Theoretical review: altered pain regulatory systems in chronic pain** *Neuroscience & Biobehavioral Reviews* **23**:877–890
19. Yarnitsky D. (2015) **Role of endogenous pain modulation in chronic pain mechanisms and treatment** *Pain* **156**:S24–S31
20. King C. D., et al. (2009) **Deficiency in endogenous modulation of prolonged heat pain in patients with irritable bowel syndrome and temporomandibular disorder** *Pain* **143**:172–178
21. Bannister K., Dickenson A. (2017) **The plasticity of descending controls in pain: translational probing** *The Journal of physiology* **595**:4159–4166
22. Tracey I. (2016) **A vulnerability to chronic pain and its interrelationship with resistance to analgesia** *Brain* **139**:869–1872
23. Jepma M., Koban L., van Doorn J., Jones M., Wager T. D. (2018) **Behavioural and neural evidence for self-reinforcing expectancy effects on pain** *Nature Human Behaviour* **2**:838–855
24. Särkkä S. (2013) **Bayesian filtering and smoothing**
25. Hoskin R., et al. (2019) **Sensitivity to pain expectations: A Bayesian model of individual differences** *Cognition* **182**:127–139
26. Kalman R. E. (1960) **A New Approach to Linear Filtering and Prediction Problems** *Transactions of the ASME- Journal of Basic Engineering* **82**:35–45
27. Sutton R. S., Barto A. G. (2018) **Reinforcement learning: an introduction**
28. Heald J. B., Lengyel M., Wolpert D. M. (2022) **Contextual inference in learning and memory** *Trends in Cognitive Sciences*
29. Piray P., Daw N. D. (2021) **A model for learning based on the joint estimation of stochasticity and volatility** *Nature Communications*
30. Mayr A., et al. (2022) **Patients with chronic pain exhibit individually unique cortical signatures of pain encoding** *Human Brain Mapping* **43**:1676–1693
31. Baliki M. N., et al. (2012) **Corticostriatal functional connectivity predicts transition to chronic back pain** *Nature neuroscience* **15**:1117–1119
32. Foss J. M., Apkarian A. V., Chialvo D. R. (2006) **Dynamics of pain: fractal dimension of temporal variability of spontaneous pain differentiates between pain states** *Journal of neurophysiology* **95**:730–736
33. Meyniel F., Dehaene S. (2017) **Brain networks for confidence weighting and hierarchical inference during probabilistic learning** *Proceedings of the National Academy of Sciences* **114**
34. Tabor A., Thacker M. A., Moseley G. L., Körding K. P. (2017) **Pain: A Statistical Account** *PLOS Computational Biology* **13**

35. Fardo F., et al. (2017) **Expectation violation and attention to pain jointly modulate neural gain in somatosensory cortex** *Neuroimage* **153**:109–121
36. Seymour B., Mancini F. (2020) **Hierarchical models of pain: Inference, information-seeking, and adaptive control** *NeuroImage* **222**
37. Büchel C., Geuter S., Sprenger C., Eippert F. (2014) **Placebo Analgesia: A Predictive Coding Perspective** *Neuron* **81**:1223–1239
38. Yoshida W., Seymour B., Koltzenburg M., Dolan R. J. (2013) **Uncertainty Increases Pain: Evidence for a Novel Mechanism of Pain Modulation Involving the Periaqueductal Gray** *Journal of Neuroscience* **33**:5638–5646
39. Anchisi D., Zanon M. (2015) **A Bayesian Perspective on Sensory and Cognitive Integration in Pain Perception and Placebo Analgesia** *PLOS ONE* **10**
40. Wiech K. (2016) **Deconstructing the sensation of pain: The influence of cognitive processes on pain perception** *Science* **354**:584–587
41. Tabor A., Burr C. (2019) **Bayesian Learning Models of Pain: A Call to Action** *Current Opinion in Behavioral Sciences* **26**:54–61
42. Colloca L., Sigaucho M., Benedetti F. (2008) **The role of learning in nocebo and placebo effects** *Pain* **136**:211–218
43. Blasini M., Corsi N., Klinger R., Colloca L. (2017) **Nocebo and pain: an overview of the psychoneurobiological mechanisms** *PAIN Reports* **2**
44. Hird E. J., Charalambous C., El-Dereby W., Jones A. K. P., Talmi D. (2019) **Boundary effects of expectation in human pain perception** *Scientific Reports* **9**
45. Brown C. A., Seymour B., El-Dereby W., Jones A. K. (2008) **Confidence in beliefs about pain predicts expectancy effects on pain perception and anticipatory processing in right anterior insula** *Pain* **139**:324–332
46. Seymour B. (2019) **Pain: a precision signal for reinforcement learning and control** *Neuron* **101**:1029–1041
47. Baliki M. N., Apkarian A. V. (2015) **Nociception, pain, negative moods, and behavior selection** *Neuron* **87**:474–491
48. Vlaeyen J. W., Crombez G., Linton S. J. (2016) **The fear-avoidance model of pain** *Pain* **157**:1588–1589
49. Strigo I. A., Carli F., Bushnell M. C. (2000) **Effect of ambient temperature on human pain and temperature perception** *Anesthesiology* **92**:699–707
50. Medoc Advanced Medical Systems (2022) **Medoc Advanced Medical Systems. TSA 2 - Advanced Thermosensory Stimulator** <https://www.medoc-web.com/tsa-2>. [Online; accessed 15-Aug-2022]. <https://www.medoc-web.com/tsa-2> (2022).
51. Lue Y.-J., Shih Y.-C., Lu Y.-M., Liu Y.-F. (2017) **Method of Limit and Method of Level for Thermal and Pain Detection Assessment** *International Journal of Physical Therapy & Rehabilitation* **3**

52. Hermans L., et al. (2016) **An Overview of Offset Analgesia and the Comparison with Conditioned Pain Modulation: A Systematic Literature Review** *Pain Physician* **19**:307–326
53. Kleiner M., et al. (2007) **What's new in psychtoolbox-3. English (US)** *Perception* **36**:1–16
54. Hollins M., Harper D., Maixner W. (2011) **Changes in pain from a repetitive thermal stimulus: the roles of adaptation and sensitization** *Pain* **152**:1583–1590
55. Virtanen P., et al. (2020) **SciPy 1.0: Fundamental Algorithms for Scientific Computing in Python** *Nature Methods* **17**:261–272
56. Stan Development Team (2019) **Stan Development Team. RStan: the R interface to Stan R package version 2.21.0. 2019.** <http://mc-stan.org/>.
57. Papaspiliopoulos O., Roberts G. O., Sköld M. (2007) **A General Framework for the Parametrization of Hierarchical Models** *Statistical Science* :59–73
58. Stan Development Team **Stan Development Team. 25.7 Reparameterization Stan User's Guide** <https://mc-stan.org/docs/stan-users-guide/reparameterization.html>. [Online; accessed 15-Aug-2022].
59. Bowling S. R., Khasawneh M. T., Kaewkuekool S., Cho B. R. (2009) **A Logistic Approximation to the Cumulative Normal Distribution** *Journal of Industrial Engineering and Management* **2**:114–27
60. Vehtari A., Gelman A., Gabry J. (2017) **Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC** *Statistics and Computing* **27**:1413–1432
61. Kruschke J. K. (2015) **Doing Bayesian data analysis: a tutorial with R, JAGS, and Stan**
62. Aylward J., et al. (2019) **Altered learning under uncertainty in unmedicated mood and anxiety disorders** *Nature Human Behaviour*
63. Ahn W.-Y., Haines N., Zhang L. (2017) **Revealing Neurocomputational Mechanisms of Reinforcement Learning and Decision-Making With the hBayesDM Package** *Computational Psychiatry* :24–57
64. Wilson R. C., Collins A. G. (2019) **Ten simple rules for the computational modeling of behavioral data** *eLife* **8**
65. White C. N., Servant M., Logan G. D. (2018) **Testing the validity of conflict drift-diffusion models for use in estimating cognitive processes: A parameter-recovery study** *Psychonomic Bulletin & Review* **25**:286–301

Editors

Reviewing Editor

Markus Ploner

Technische Universität München, Munich, Germany

Senior Editor

Christian Büchel

University Medical Center Hamburg-Eppendorf, Hamburg, Germany

Reviewer #1 (Public Review):**Summary:**

This study examined the role of statistical learning in pain perception, suggesting that individuals' expectations about a sequence of events influence their perception of pain intensity. They incorporated the components of volatility and stochasticity into their experimental design and asked participants ($n = 27$) to rate the pain intensity, their prediction, and their confidence level. They compared two different inference strategies: Bayesian inference vs. heuristic-employing Kalman filters and model-free reinforcement learning. They showed that the expectation-weighted Kalman filter best explained the temporal pattern of participants' ratings. These results provide evidence for a Bayesian inference perspective on pain, supported by a computational model that elucidates the underlying process.

Strengths:

- Their experimental design included a wide range of input intensities and the levels of volatility and stochasticity. With elaborated computational models, they provide solid evidence that statistical learning shapes pain.

Weaknesses:

- Relevance to clinical pain: While the authors underscore the relevance of their findings to chronic pain, they did not include data pertaining to clinical pain.

<https://doi.org/10.7554/eLife.90634.2.sa1>

Reviewer #3 (Public Review):

The study investigated how statistical aspects of temperature sequences, such as manipulations of stochasticity (i.e., randomness of a sequence) and volatility (i.e., speed at which a sequence unfolded) influenced pain perception. Using an innovative stimulation paradigm and computational modelling of perceptual variables, this study demonstrated that perception is weighted by expectations. Overall, the findings support the conclusion that pain perception is mediated by expectations in a Bayesian manner. The provision of additional details during the review process strengthens the reliability of this conclusion. The methods presented offer tools and frameworks for further research in pain perception and can be extended to investigations into chronic pain processes.

<https://doi.org/10.7554/eLife.90634.2.sa0>

Author response:

The following is the authors' response to the original reviews.

eLife assessment

This study presents a valuable insight into a computational mechanism of pain perception. The evidence supporting the authors' claims is solid, although the inclusion of 1) more diverse candidate computational models, 2) more systematic analysis of the temporal regularity effects on the model fit, and 3) tests on clinical samples would have strengthened the study. The work will be of interest to pain researchers working on computational models and cognitive mechanisms of pain in a Bayesian framework.

Thank you very much again for considering the manuscript and judging it as a valuable contribution to understanding mechanisms of pain perception. We recognise the above-mentioned points of improvement and elaborate on them in the initial response to the reviewers.

Response to the reviewers

Reviewer 1:

Reviewer Comment 1.1 — Selection of candidate computational models: While the paper juxtaposes the simple model-free RL model against a Kalman Filter model in the context of pain perception, the rationale behind this choice remains ambiguous. It prompts the question: could other RL-based models, such as model-based RL or hierarchical RL, offer additional insights? A more detailed explanation of their computational model selection would provide greater clarity and depth to the study.

Initial reply: Thank you for this point. Our models were selected a-priori, following the modelling strategy from Jepma et al. (2018) and hence considered the same set of core models for clear extension of the analysis to our non-cue paradigm. The key question for us was whether expectations were used to weight the behavioural estimates, so our main interest was to compare expectation vs non-expectation weighted models.

Model-based and hierarchical RL are very broad terms that can be used to refer to many different models, and we are not clear about which specific models the reviewer is referring to. Our Bayesian models are generative models, i.e. they learn the generative statistics of the environment (which is characterised by inherent stochasticity and volatility) and hence operate model-based analyses of the stimulus dynamics. In our case, this happened hierarchically and it was combined with a simple RL rule.

Revised reply: We clarified our modelling choices in the "Modelling strategy" subsection of the results section.

Reviewer Comment 1.2 — Effects of varying levels of volatility and stochasticity: The study commendably integrates varying levels of volatility and stochasticity into its experimental design. However, the depth of analysis concerning the effects of these variables on model fit appears shallow. A looming concern is whether the superior performance of the expectation-weighted Kalman Filter model might be a natural outcome of the experimental design. While the non-significant difference between eKF and eRL for the high stochasticity condition somewhat alleviates this concern, it raises another query: Would a more granular analysis of volatility and stochasticity effects reveal fine-grained model fit patterns?

Initial reply: We are sorry that the reviewer finds shallow "the depth of analysis concerning the effects of these variables on model fit". We are not sure which analysis the reviewer has in mind when suggesting a "more granular analysis of volatility and stochasticity effects" to "reveal fine-grained model fit patterns". Therefore, we find it difficult to improve our manuscript in this regard. We are happy to add analyses to our paper but we would be grateful for some specific pointers. We have already provided:

- Analysis of model-naïve performance across different levels of stochasticity and volatility (section 2.3, figure 3, supplementary information section 1.1 and tables S1-2)
- Model fitting for each stochasticity/volatility condition (section 2.4.1, figure 4, supplementary table S5)

- Group-level and individual-level differences of each model parameter across stochasticity/volatility conditions (supplementary information section 7, figures S4-S5).
- Effect of confidence on scaling factor for each stochasticity/volatility condition (figure 5)

*Reviewer Comment 1.3 — Rating instruction: According to Fig. 1A, participants were prompted to rate their responses to the question, “How much pain DID you just feel?” and to specify their confidence level regarding their pain. It is difficult for me to understand the meaning of confidence in this context, given that they were asked to report their *subjective* feelings. It might have been better to query participants about perceived stimulus intensity levels. This perspective is seemingly echoed in lines 100-101, “the primary aim of the experiment was to determine whether the expectations participants hold about the sequence inform their perceptual beliefs about the intensity of the stimuli.”*

Initial reply: Thank you for raising this question, which allows us to clarify our paradigm. On half of the trials, participants were asked to report the perceived intensity of the previous stimulus; on the remaining trials, participants were requested to predict the intensity of the next stimulus. Therefore, we did query “participants about perceived stimulus intensity levels”, as described at lines 49-55, 296-303, and depicted in figure 1.

The confidence refers to the level of confidence that participants have regarding their rating - how sure they are. This is done in addition to their perceived stimulus intensity and it has been used in a large body of previous studies in any sensory modality.

Reviewer Comment 1.4 — Relevance to clinical pain: While the authors underscore the relevance of their findings to chronic pain, they did not include data pertaining to clinical pain. Notably, their initial preprint seemed to encompass data from a clinical sample (<https://www.medrxiv.org/content/10.1101/2023.03.23.23287656v1>), which, for reasons unexplained, has been omitted in the current version. Clarification on this discrepancy would be instrumental in discerning the true relevance of the study's findings to clinical pain scenarios.

Initial reply: The preprint that the Reviewer is referring to was an older version of the manuscript in which we combined two different experiments, which were initially born as separate studies: the one that we submitted to eLife (done in the lab, with noxious stimuli in healthy participants) and an online study with a different statistical learning paradigm (without noxious stimuli, in chronic back pain participants). Unfortunately, the paradigms were different and not directly comparable. Indeed, following submission to a different journal, the manuscript was criticised for this reason. We therefore split the paper in two, and submitted the first study to eLife. We are now planning to perform the same lab-based experiment with noxious stimuli on chronic back pain participants. Progress on this front has been slowed down by the fact that I (Flavia Mancini) am on maternity leave, but it remains top priority once back to work.

Reviewer Comment 1.5 — Paper organization: The paper's organization appears a little bit weird, possibly due to the removal of significant content from their initial preprint. Sections 2.12.2 and 2.4 seem more suitable for the Methods section, while 2.3 and 2.4.1 are the only parts that present results. In addition, enhancing clarity through graphical diagrams, especially for the experimental design and computational models, would be quite beneficial. A reference point could be Fig. 1 and Fig. 5 from Jepma et al. (2018), which similarly explored RL and KF models.

Initial reply: Thank you for these suggestions. We will consider restructuring the paper in the revised version.

Revised reply: We restructured introduction, results and parts of the methods. We followed the reviewer's suggestion regarding enhancing clarity through graphical diagrams. We have visualised the experimental design in Figure 1D. Furthermore, we have visualised the two main computational models (eRL and eKF) in Figure 2, following from Jepma et al. (2018). As a result, we have updated the notation in Section 4.4 to be clearer and consistent with the graphical representation (rename the variable referring to observed thermal input from O_t to N_t).

Reviewer Comment 1.6 — In lines 99-100, the statement “following the work by [23]” would be more helpful if it included a concise summary of the main concepts from the referenced work.

- It would be helpful to have descriptions of the conditions that Figure 1C is elaborating on.

- In line 364, the “ $N_{\{t\}}$ ” in the sentence “The observation on trial t , $N_{\{t\}}$ ”, should be $O_{\{t\}}$.

Initial reply: Thank you for spotting these and for providing the suggestions. We will include the correction in the revised version.

Revised reply: We have added the following regarding the lines 99-100:

“We build on the work by [23], who show that pain perception is strongly influenced by expectations as defined by a cue that predicts high or low pain. In contrast to the cue-paradigm from [23], the primary aim of our experiment was to determine whether the expectations participants hold about the sequence itself inform their perceptual beliefs about the intensity of the stimuli.”

See comment in the previous reply, regarding the notation change from O_t to N_t .

Reviewer 2:

Reviewer Comment 2.1 — This is a highly interesting and novel finding with potential implications for the understanding and treatment of chronic pain where pain regulation is deficient. The paradigm is clear, the analysis is state-of-the-art, the results are convincing, and the interpretation is adequate.

Initial reply: Thank you very much for these positive comments.

Reviewer 3:

Summary:

I am pleased to have had the opportunity to review this manuscript, which investigated the role of statistical learning in the modulation of pain perception. In short, the study showed that statistical aspects of temperature sequences, with respect to specific manipulations of stochasticity (i.e., randomness of a sequence) and volatility (i.e., speed at which a sequence unfolded) influenced pain perception. Computational modelling of perceptual variables (i.e., multi-dimensional ratings of perceived or predicted stimuli) indicated that models of perception weighted by expectations were the best explanation for the data. My comments below are not intended to undermine or question the quality of this research. Rather, they are offered with the intention of enhancing what is already a significant contribution to the pain neuroscience field. Below, I highlight the strengths

and weaknesses of the manuscript and offer suggestions for incorporating additional methodological details.

Strengths:

The manuscript is articulate, coherent, and skilfully written, making it accessible and engaging.

- The innovative stimulation paradigm enables the exploration of expectancy effects on perception without depending on external cues, lending a unique angle to the research.

- By including participants' ratings of both perceptual aspects and their confidence in what they perceived or predicted, the study provides an additional layer of information to the understanding of perceptual decision-making. This information was thoughtfully incorporated into the modelling, enabling the investigation of how confidence influences learning.

- The computational modelling techniques utilised here are methodologically robust. I commend the authors for their attention to model and parameter recovery, a facet often neglected in previous computational neuroscience studies.

- The well-chosen citations not only reflect a clear grasp of the current research landscape but also contribute thoughtfully to ongoing discussions within the field of pain neuroscience.

Initial reply: We are really grateful for reviewer's insightful comments and for providing useful guidance regarding our methodology. We are also thankful for highlighting the strengths of our manuscript. Below we respond to individual weakness mentioned in the reviews report.

Reviewer Comment 3.1 — In Figure 1, panel C, the authors illustrate the stimulation intensity, perceived intensity, and prediction intensity on the same scale, facilitating a more direct comparison. It appears that the stimulation intensity has been mathematically transformed to fit a scale from 0 to 100, aligning it with the intensity ratings corresponding to either past or future stimuli. Given that the pain threshold is specifically marked at 50 on this scale, one could logically infer that all ratings falling below this value should be deemed non-painful. However, I find myself uncertain about this interpretation, especially in relation to the term "arbitrary units" used in the figure. I would greatly appreciate clarification on how to accurately interpret these units, as well as an explanation of the relationship between these values and the definition of pain threshold in this experiment.

Initial reply: Indeed, as detailed in the Methods section 4.3, the stimulation intensity was originally transformed from the 1-13 scale to 0-100 scale to match the scales in the participant response screens.

Following the method used to establish the pain threshold, we set the stimulus intensity of 7 as the threshold on the original 1-13 scale. However, during the rating part of the experiment, several of the participants never or very rarely selected a value above 50 (their individually defined pain threshold), despite previously indicating a moment during pain threshold procedure when a stimulus becomes painful. This then results in the re-scaled intensity values as well the perception rating, both on the same 0-100 scale of arbitrary units, to never go above the pain threshold. Please see all participant ratings and inputs in the Figure below. We see that it would be more illustrative to re-plot Figure 1 with a different exemplary participant, whose ratings go above the pain threshold, perhaps with an input intensity on

the 1-13 scale on the additional right-hand-side y-axis. We will add this in the revised version as well as highlight the fact above.

Importantly, while values below 50 are deemed non-painful by participants, the thermal stimulation still activates C-fibres involved in nociception, and we would argue that the modelling framework and analysis still applies in this case.

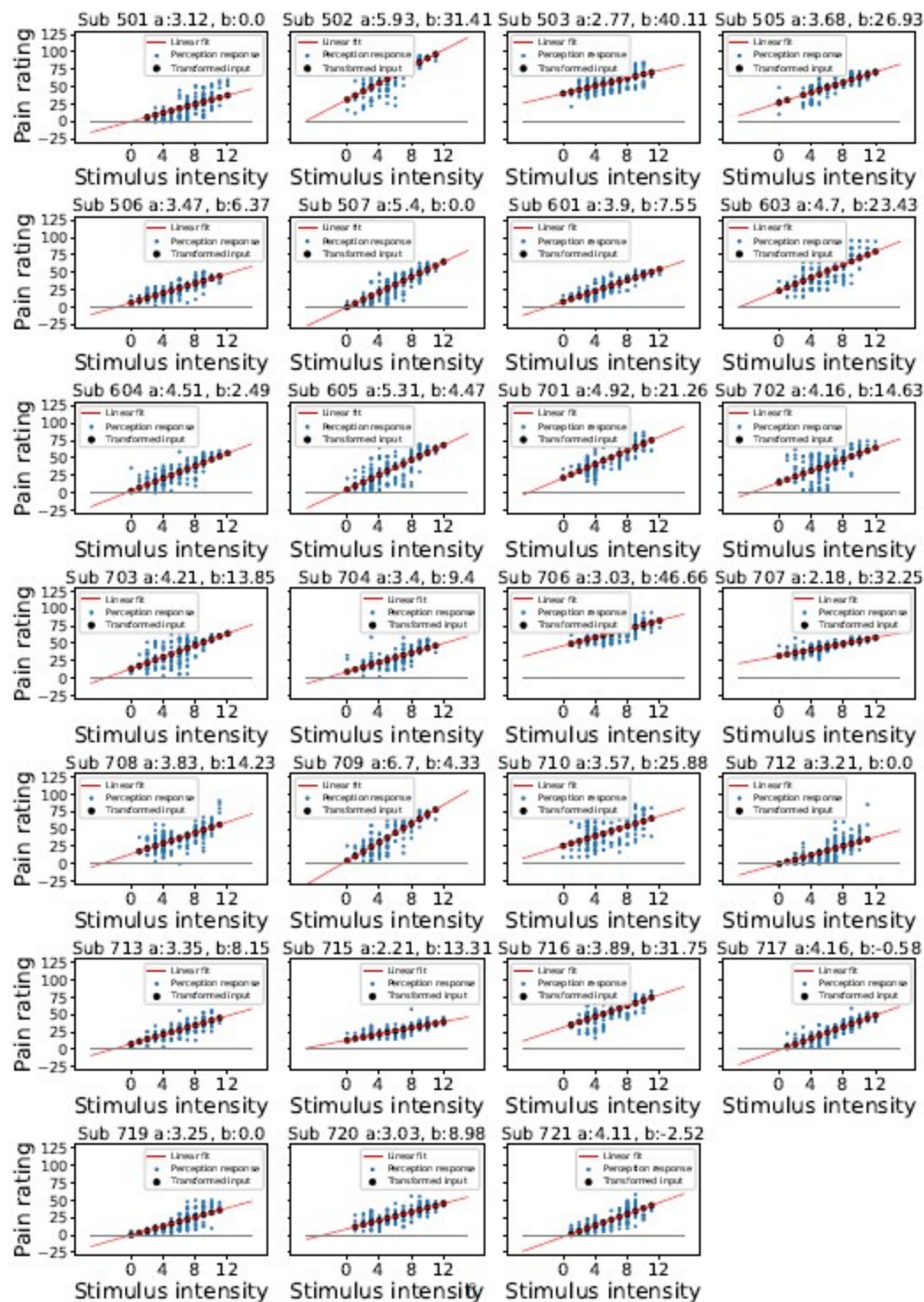
Revised reply: We re-plotted Figure 1E-F with a different exemplary participant, whose rating go above the pain threshold. We also included all participant pain perception and prediction ratings, noxious input sequences and confidence ratings in the supplement in Figures S1-S3.

Reviewer Comment 3.2 — The method of generating fluctuations in stimulation temperatures, along with the handling of perceptual uncertainty in modelling, requires further elucidation. The current models appear to presume that participants perceive each stimulus accurately, introducing noise only at the response stage. This assumption may fail to capture the inherent uncertainty in the perception of each stimulus intensity, especially when differences in consecutive temperatures are as minimal as 1°C.

Initial reply: We agree with the reviewer that there are multiple sources of uncertainty involved in the process of rating the intensity of thermal stimuli - including the perception uncertainty. In order to include an account of inaccurate perception, one would have to consider different sources that contribute to this, which there may be many. In our approach, we consider one, which is captured in the expectation weighted model, more clearly exemplified in the expectation-weighted Kalman-Filter model (eKF). The model assumes participants perception of input as an imperfect indicator of the true level of pain. In this case, it turns out that perception is corrupted as a result of the expectation participants hold about the upcoming stimuli. The extent of this effect is partly governed by a subjective level of noise ϵ , which may also subsume other sources of uncertainty beyond the expectation effect. Moreover, the response noise ξ , could also subsume any other unexplained sources of noise.

Author response image 1.

Stimulus intensity transformation



Revised reply: We clarified our modelling choices in the "2.2 Modelling strategy" subsection.

Reviewer Comment 3.3 — A key conclusion drawn is that eKF is a better model than eRL. However, a closer examination of the results reveals that the two models behave very similarly, and it is not clear that they can be readily distinguished based on model recovery and model comparison results.

Initial reply: While, the eKF appears to rank higher than the eRL in terms of LOOIC and sigma effects, we don't wish to make sweeping statements regarding significance of differences between eRL and eKF, but merely point to the trend in the data. We shall make this clearer in the revised version of the manuscript. However, the most important result is that the models involving expectation-weighting are arguably better capturing the data.

Revised reply: We elaborated on the significance statements in the "Modelling Results" subsection:

- We considered at least a 2 sigma effect as indication of a significant difference. In each condition, the expectation weighted models (eKF and eRL) provided better fit than models without this element (KF and RL; approx. 2-4 sigma difference, as reported in Figure 5A-D). This suggests that regardless of the levels of volatility and stochasticity, participants still weigh perception of the stimuli with their expectation.

and in the first paragraph of the Discussion:

- When varying different levels of inherent uncertainty in the sequences of stimuli (stochasticity and volatility), the expectation and confidence weighted models fitted the data better than models weighted for confidence but not for expectations (Figure 5A-D). The expectation-weighted bayesian (KF) model offered a better fit than the expectation-weighted, model-free RL model, although in conditions of high stochasticity this difference was short of significance. Overall, this suggests that participants' expectations play a significant role in the perception of sequences of noxious stimuli.

We are aware of the limitations and lack of clear guidance regarding using sigma effects to establish significance (as per reviewer's suggestion: <https://discourse.mc-stan.org/t/loo-comparison-in-referenceto-standard-error/4009>). Here we decided to use the above-mentioned threshold of 2-sigma as an indication of significance, but note the potential limitations of the inferences - especially when distinguishing between eRL/eKF models.

Reviewer Comment 3.4 — Regarding model recovery, the distinction between the eKF and eRL models seems blurred. When the simulation is based on the eKF, there is no ability to distinguish whether either eKF or eRL is better. When the simulation is based on the eRL, the eRL appears to be the best model, but the difference with eKF is small. This raises a few more questions. What is the range of the parameters used for the simulations?

Initial reply: We agree that the distinction between eKF and eRL in the model recovery is not that clean-cut, which may in turn point to the similarity between the two models. To simulate the data for the model and parameter recovery analysis, we used the group means and variances estimated on the participant data to sample individual parameter values.

Reviewer Comment 3.5 — Is it possible that either eRL or eKF are best when different parameters are simulated? Additionally, increasing the number of simulations to at least 100 could provide more convincing model recovery results.

Initial reply: It could be a possibility, but would require further investigation and comparison of fits for different bins/ranges of parameters to see if there is any consistent advantage of one model over another in each bin. We will consider adding this analysis, and provide an additional 50 simulations to paint a more convincing picture.

Revised reply: We increased the number of simulations per model pair to ≈ 100 (after rejecting fits based on diagnostics criteria - E-BFMI and divergent transitions) and updated the confusion matrix (Table S4). Although the confusion between eRL and eKF remains, the

model recovery shows good distinction between expectation weighted vs non-expectation weighted (and Random) models, which supports our main conclusion in the paper.

Reviewer Comment 3.6 — Regarding model comparison, the authors reported that “the expectation-weighted KF model offered a better fit than the eRL, although in conditions of high stochasticity, this difference was short of significance against the eRL model.” This interpretation is based on a significance test that hinges on the ratio between the ELPD and the surrounding standard error (SE). Unfortunately, there’s no agreed-upon threshold of SEs that determines significance, but a general guideline is to consider “several SEs,” with a higher number typically viewed as more robust. However, the text lacks clarity regarding the specific number of SEs applied in this test. At a cursory glance, it appears that the authors may have employed 2 SEs in their interpretation, while only depicting 1 SE in Figure 4.

Initial reply: Indeed, we considered 2 sigma effect as a threshold, however we recognise that there is no agreed-upon threshold, and shall make this and our interpretation clearer regarding the trend in the data, in the revision.

Revised reply: We clarify this further, as per our revised response to Comment 3.3 above. We have also added the following statement in section 4.5.1 (Methods, Model comparison): “There’s no agreed-upon threshold of SEs that determines significance, but the higher the sigma difference, the more robust is the effect.”

Reviewer Comment 3.7 — With respect to parameter recovery, a few additional details could be included for completeness. Specifically, while the range of the learning rate is understandably confined between 0 and 1, the range of other simulated parameters, particularly those without clear boundaries, remains ambiguous. Including scatter plots with the simulated parameters on the xaxis and the recovered parameters on the y-axis would effectively convey this missing information.

Furthermore, it would be beneficial for the authors to clarify whether the same priors were used for both the modelling results presented in the main paper and the parameter recovery presented in the supplementary material.

Initial reply: Thanks for this comment and for the suggestions. To simulate the data for the model and parameter recovery analysis, we used the group means and variances estimated on the participant data to sample individual parameter values. The priors on the group and individual-level parameters in the recovery analysis were the same as in the fitting procedure. We will include the requested scatter plots in the next iteration of the manuscript.

Revised reply: We included parameter recovery scatter plots for each model and parameter in the Supplement Figures S7-S11.

Reviewer Comment 3.8 — While the reliance on R-hat values for convergence in model fitting is standard, a more comprehensive assessment could include estimates of the effective sample size (bulk ESS and/or tail ESS) and the Estimated Bayesian Fraction of Missing Information (EBFMI), to show efficient sampling across the distribution. Consideration of divergences, if any, would further enhance the reliability of the results.

Initial reply: Thank you very much for this suggestion, we will aim to include these measures in the revised version.

Revised reply: We have considered the suggested diagnostics and include bulk and tail ESS values for each condition, model, parameter in the Supplement Tables S6-S9. We also report

number of chain with low E-BFMI (0), number of divergent transitions (0) and the E-BFMI values per chain in Table S10.

Reviewer Comment 3.9 — The authors write: “Going beyond conditioning paradigms based in cuing of pain outcomes, our findings offer a more accurate description of endogenous pain regulation.” Unfortunately, this statement isn’t substantiated by the results. The authors did not engage in a direct comparison between conditioning and sequence-based paradigms. Moreover, even if such a comparison had been made, it remains unclear what would constitute the gold standard for quantifying “endogenous pain regulation.”

Initial reply: This is valid point, indeed we do not compare paradigms in our study, and will remove this statement in the future version.

Revised reply: We have removed this statement from the revised version.

Reviewer Comment 3.10 — In relation to the comment on model comparison in my public review, I believe the following link may provide further insight and clarify the basis for my observation. It discusses the use of standard error in model comparison and may be useful for the authors in addressing this particular point: <https://discourse.mc-stan.org/t/loo-comparison-in-referenceto-standard-error/4009>

Initial reply: Thank you for this suggestion, we will consider the forum discussion in our manuscript.

<https://doi.org/10.7554/eLife.90634.2.sa3>