

1 We thank the three anonymous reviewers for their insightful and largely positive comments!

2 **Responses to Reviewer 1** Thanks again for your thoughtful comments.

3 **Lemmas 3.1 and 3.2.** Regarding Lemma 3.1, the novelty is the existence of the function $q(x)$. If $q(x) = 1$, this is well
4 known in the literature, as the reviewer pointed out. We will make this point clear in the revision. Regarding Lemma
5 3.2, while we do not see this as a key contribution of the paper, it is not clear to us whether it follows from the results of
6 Smola and Schoelkopf. This is because they consider greedy selection of a subset of training data, while we consider
7 selection of points from the entire domain. If you know any existing work that explicitly states Lemmas 3.2, could you
8 please let us know when you update your review? We will cite it and remove our proof in Appendix.

9 **Lemma 4.1.** Thank you for the suggestion, which we follow in the revision.

10 **Proof of Proposition 4.2.** Thank you very much for pointing this out. We agree that our argument of deriving (v) was
11 flawed. We make the following correction, which we believe makes Proposition 4.2 still valid. (We will correct the proof
12 of Prop. C.2 in a similar way.) We start from (iii), which can be stated as that there exists $n_0 \in \mathbb{N}$ such that $k_{X_n}(x, x) \leq$
13 $\exp(-(c_1/c_2)n^{1/d})$ holds for all $n \geq n_0$. Now, define $c_4 > 0$ as a constant such that $c_4 \exp(-(c_1/c_2)n_0^{1/d}) = c_3$, and
14 let $c_5 := \max(c_4, 1)$. Then, for $n < n_0$ we have $c_5 \exp(-(c_1/c_2)n^{1/d}) \geq c_4 \exp(-(c_1/c_2)n^{1/d}) \geq c_3 \geq k_{X_n}(x, x)$,
15 where the last inequality follows from (iv). For $n \geq n_0$, we have $c_5 \exp(-(c_1/c_2)n^{1/d}) \geq \exp(-(c_1/c_2)n^{1/d}) \geq$
16 $k_{X_n}(x, x)$. Therefore we conclude that $k_{X_n}(x, x) \leq c_5 \exp(-(c_1/c_2)n^{1/d})$ holds for all $n \in \mathbb{N}$ and $x \in \Omega$.

17 **Responses to Reviewer 2** Thanks again for your insightful comments.

18 **The estimator in line 120.** This is a very good point, and thanks for pointing it out. We agree that the current
19 presentation is confusing, and will make a correction. The quadrature estimator suggested in [8] (and used for WSABI-
20 M [13]) can be described as $\int \mathbb{E}_{\hat{g}} T(\hat{g}(x)) \pi(x)$, where $\hat{g} \sim \mathcal{GP}(m_{g, X_n}, k_{X_n})$ is the posterior GP. On the other hand, the
21 estimator in line 120 is $\int T(m_{g, X_n}(x)) \pi(x)$ and used by WSABI-L [13]. As we describe below, these two estimators
22 are *both consistent with the same convergence rates*, and all theoretical guarantees obtained in the paper are applicable
23 to the estimator $\int \mathbb{E}_{\hat{g}} T(\hat{g}(x)) \pi(x)$ as well. Intuitively, this is because the posterior $\hat{g} \sim \mathcal{GP}(m_{g, X_n}, k_{X_n})$ contracts
24 around the posterior mean m_{g, X_n} as n increases, and $\mathbb{E}_{\hat{g}} T(\hat{g})$ and $T(m_{g, X_n})$ get similar. We note however that for a
25 finite n , it is not clear which estimator is “better,” and this is an interesting topic for future research.

26 We sketch here that for the estimator $\int \mathbb{E}_{\hat{g}} T(\hat{g}(x)) \pi(x)$, the essentially same upper bound as Proposition 2.1 holds, under
27 an additional condition that $\mathbb{E}_{\hat{g}} (T'(|g(x)| + |\hat{g}(x)|))^2 < C$ holds for all $x \in \Omega$ and $n \in \mathbb{N}$ for some $C > 0$ (which can be
28 shown to be satisfied for transformations T mentioned in our paper). By Taylor’s theorem, there exists $\alpha_{x, X_n, \hat{g}} \in [0, 1]$
29 such that for $y_{x, X_n, \hat{g}} := g(x) + \alpha_{x, X_n, \hat{g}}(\hat{g}(x) - g(x))$ we have $T(\hat{g}(x)) = T(g(x)) + T'(y_{x, X_n, \hat{g}})(\hat{g}(x) - g(x))$.
30 Therefore $(\mathbb{E}_{\hat{g}}[T(\hat{g}(x))] - T(g(x)))^2 = (\mathbb{E}_{\hat{g}}[T'(y_{x, X_n, \hat{g}})(\hat{g}(x) - g(x))])^2 \leq \mathbb{E}_{\hat{g}}[(T'(y_{x, X_n, \hat{g}}))^2] \mathbb{E}_{\hat{g}}[(\hat{g}(x) - g(x))^2] \leq$
31 $C \mathbb{E}_{\hat{g}}[(\hat{g}(x) - g(x))^2]$, where the last inequality follows from $|y_{x, X_n, \hat{g}}| \leq |g(x)| + |\hat{g}(x)|$ and the above assumption.
32 Moreover, $\mathbb{E}_{\hat{g}}[(\hat{g}(x) - g(x))^2] \leq 2\mathbb{E}_{\hat{g}}[(\hat{g}(x) - m_{g, X_n}(x))^2] + 2(m_{g, X_n}(x) - g(x))^2 \leq 2k_{X_n}(x, x) + 2\|\hat{g}\|_{\mathcal{H}_k}^2 k_{X_n}(x, x)$,

33 where the last inequality follows from Eq. (10). Thus, $|T(g(x)) - \mathbb{E}_{\hat{g}}[T(\hat{g}(x))]| \leq \sqrt{2C(1 + \|\hat{g}\|_{\mathcal{H}_k}^2)} \sqrt{k_{X_n}(x, x)}$.
34 Following the argument in line 436, the essentially same bound as Prop. 2.1 can be obtained (with a different constant).

35 **The generic form of acquisition functions (Eq. 4).** We defined Eq. (4) so that the class of acquisition functions to
36 which our convergence guarantees are applicable becomes as large as possible. One positive side of this generality
37 is that it enables practitioners to design a new acquisition that results in a consistent algorithm; its derivation can be
38 quite different from those of the existing acquisition functions (which are often done by approximating an intractable
39 posterior covariance function of the transformed integrand). In this sense, we think that every special case of (4) does
40 not have to have an interpretation as an approximate posterior covariance. We nevertheless agree that Eq. 4 is rather
41 abstract. We will include a discussion about the roles of the components in Eq. (4).

42 **Responses to Reviewer 3** Thanks again for encouraging comments and insightful suggestions.

43 **BQ v.s. MC.** We will make the tone of the comparison more neutral, including a discussion of the dependence of the
44 dimensionality. We hope that you agree, however, that the existence of a convergence guarantee for MCMC (which
45 didn’t exist for adaptive BQ so far) is a key contributor to MCMC’s popularity.

46 **Ergodicity / detailed balance and the relation to weak adaptivity.** We will modify the presentation so that the
47 intuitive discussion on the connection to the detailed balance and ergodicity becomes minimal. We will also add an
48 intuitive explanation about the weak adaptivity condition immediately after it is introduced.

49 **Proof map / Appendix.** We will provide a diagram in Appendix that describes the relationships between the various
50 auxiliary results and how they yield the main results. We will also add a high level overview of the proof plan and
51 techniques used.