

Tight Runtime Bounds for Static Unary Unbiased Evolutionary Algorithms on Linear Functions

Carola Doerr^{1†}, Duri Andrea Janett^{1,2†}, Johannes Lengler^{2*†}

¹*LIP6, CNRS, Sorbonne Université, Paris, France.

²Department of Computer Science, ETH Zürich, Zürich, Switzerland.

*Corresponding author(s). E-mail(s): johannes.lengler@inf.ethz.ch;

†All authors contributed equally to this work.

Abstract

In a seminal paper in 2013, Witt showed that the (1+1) Evolutionary Algorithm with standard bit mutation needs time $(1 + o(1))n \ln n / p_1$ to find the optimum of any linear function, as long as the probability p_1 to flip exactly one bit is $\Theta(1)$. In this paper we investigate how this result generalizes if standard bit mutation is replaced by an arbitrary unbiased mutation operator. This situation is notably different, since the stochastic domination argument used for the lower bound by Witt no longer holds. In particular, starting closer to the optimum is not necessarily an advantage, and OneMax is no longer the easiest function for arbitrary starting positions.

Nevertheless, we show that Witt’s result carries over if p_1 is not too small, with different constraints for upper and lower bounds, and if the number of flipped bits has bounded expectation χ . Notably, this includes some of the heavy-tail mutation operators used in fast genetic algorithms, but not all of them. We also give examples showing that algorithms with unbounded χ have qualitatively different trajectories close to the optimum.

Keywords: Runtime analysis, Theory of Evolutionary Computation, Mutation Operators

1 Introduction

One of the most crucial ingredients of evolutionary algorithms is the *mutation operator*, i.e., the procedure that describes how to generate offspring from a single parent. On the hypercube $\{0, 1\}^n$, for a long time the undisputed default was to use *standard bit mutation*, which flips each bit of the parent independently with the same probability. However, this convention has been challenged in the last years; for example via the *fast* mutation operators [DLMN17], for which the number of flipped bits follows a heavy-tailed distribution. The advantages of using heavy-tailed distributions are rather impressive [DN21]. They are slightly worse for hillclimbing, but the expected runtime deteriorates only by a constant factor that can be chosen close to one. However, they are massively better at escaping local optima. While it takes $e^{\Omega(k \ln k)}$ steps to make a jump of size k with standard bit mutation of mutation rate $\Theta(1/n)$, it only takes $k^{O(1)}$ steps with fast mutation operators. Consequently, they are faster on landscapes with local optima, like the JUMP function [DLMN17, AD20] and its generalizations [BBD22, AN21], or random MAX-3SAT instances [ABD22].¹ Heavy-tailed distributions can also help on unimodal landscapes like ONEMAX. For example, the $(1 + (\lambda, \lambda))$ GA [DDE15] was shown to achieve linear expected runtime [ABD22] when equipped with fast mutation operators, which is asymptotically best possible.

¹When the required jump size k is known in advance, then choosing the mutation rate to be k/n is optimal, as shown in [DLMN17]. The advantage of the fast mutation operator is that k does not need to be known.

Other benchmarks on which fast mutation operators or other unbiased mutation operators than standard bit mutation have been found to be useful include theoretical benchmarks like LEADINGONES [YDB19] and TwOMax [FQW18], network problems like maximum cut [FGQW18, FQW18, QGWF21], minimum vertex cover [FQW18, Buz22], maximum independent set [YWDB20], maximum flow [MB17] and SAT [SCP⁺21], landscape classes like submodular functions [FGQW18, QGWF21] and random NK-landscapes [YWDB20], multi-objective settings [DZ21, DQ22, DHP22] and other problems like subset selection [WQT18], the N-queens problem [YWDB20], the symmetric mutual information problem [FGQW18, QGWF21] and many more [YWDB20, NSN⁺22, NXN22, KNSH21, DGB22, ABD21]. Such mutation operators are integrated into standard benchmarking tools like the IOHprofiler [DWY⁺18] and Nevergrad [BDM⁺21], and they have been used as building blocks for more sophisticated algorithms [COY21a, COY21b, DR22, NAN22, PCS22].

The large success of non-standard mutation operators raises the desire to analyze which operators are (provably) optimal for a given problem setting. Such questions can be answered in the black-box complexity framework proposed in [DJW06] (see [Doe20] for a survey on the role of black-box complexity for evolutionary computation). Particularly interesting for the study of mutation operators is the unary unbiased black-box complexity model defined in [LW12]. Unary unbiased black-box algorithms create solution candidates by sampling uniformly at random or by selecting one previously evaluated point x and a search radius r (both possibly random) and then sampling the solution candidate uniformly at random among all points at Hamming distance r from x . The unary unbiased black-box complexity of a collection $\mathcal{F}$ of functions is then the best (over all unary unbiased algorithms) worst-case (over all problem instances in $\mathcal{F}$) expected runtime. The study of unary unbiased black-box complexities has led to important insights into the limitation of mutation-based algorithms [DJK⁺11, DW12, LW12, DKLW13, DD14, DDK15, LS19, DDY20], which were exploited for the design of faster algorithms such as the $(1 + (\lambda, \lambda))$ GA in [DDE15].

For ONEMax, a tight bound for the unary unbiased black-box complexity was proven in [DDY20]. It was shown there that the *drift-maximizing* algorithm that at every step chooses the mutation operator that maximizes the expected progress achieves asymptotically optimal expected runtime, up to small lower order terms. Zooming further into this problem for concrete dimensions, Buskulić and Doerr [BD21a] showed that slightly better performance can be achieved by increasing the mutation rates, i.e., by implementing a more risky strategy that, at several stages that are sufficiently far away from the optimum, flips more bits (in the hope of making more progress and at the cost of a smaller success probability). The approach developed by [BD21a] was later extended in [BD20] to compute the optimal mutation rates for the $(1 + 1)$ -EA and the $(1 + \lambda)$ -EA optimizing ONEMax. The best *static* unary unbiased mutation operator for the $(1 + \lambda)$ -EA for a number of different combinations of n and λ was numerically approximated in [BD21b]. In particular, it was shown there that the optimal mutation operators are none of the standard choices that are typically used in evolutionary algorithms. These results demonstrate that even for the optimization of ONEMax our understanding of optimal mutation operators is rather limited, both in the static and in the dynamic case.

Our Results: We aim to extend in this work the above-mentioned results to the optimization of a larger class of functions. The first natural extension of ONEMax are linear functions, so we primarily focus on these. Our particular aim is to derive tight bounds for the expected runtime of the $(1 + 1)$ -EA equipped with an arbitrary unary unbiased mutation operator.

To express our main result, we briefly recall from [Doe20] that every unbiased mutation operator can be described by a sequence of $n + 1$ probabilities $p_0, p_1, \dots, p_n$ that sum up to one. We thus identify the mutation operator with the sequence $\mathcal{D} = (p_0, p_1, \dots, p_n)$ and write $(1 + 1)$ -EA $_{\mathcal{D}}$ for the $(1 + 1)$ -EA that generates its solution candidates using the mutation operator $\text{mut}_{\mathcal{D}}$ that first draws an index $i \in [0, n]$ according to the probabilities (i.e., it picks i with probability p_i), and then flips a uniformly random set of exactly i bits. Every $(1 + 1)$ -EA equipped with an arbitrary but static unary unbiased mutation operator can be expressed as a $(1 + 1)$ -EA $_{\mathcal{D}}$. We show the following.

Theorem 1. *Consider the $(1 + 1)$ -EA $_{\mathcal{D}}$ for a distribution $\mathcal{D} = (p_0, p_1, \dots, p_n)$ with mean χ . If $p_1 = \Theta(1)$ and $\chi = O(1)$, then the expected runtime on any linear function on $\{0, 1\}^n$ with non-zero weights is*

$$(1 \pm o(1)) \frac{1}{p_1} \cdot n \ln n. \quad (1)$$

More precise versions of Theorem 1 will be presented in Corollary 10 and Theorem 14. In particular, we will show that the lower bound holds for any function with unique global optimum if $p_{n-1} = o(p_1)$. Moreover, the conditions on p_1 and χ in Theorem 1 can be slightly relaxed. We show that the expected runtime remains unchanged if $\chi^3 p_1^{-2} (1 - p_0)^{-1} = o(\ln n / \ln \ln n)$, which is probably not tight. However, we also show that the condition is not superfluous either. If p_1 becomes too small, or χ becomes too large, then the behavior of the algorithm starts to change, see Section 3.1.

Theorem 1 can be seen as a generalization of Witt's seminal work [Wit13] on linear functions, where he showed that the expected runtime of the $(1 + 1)$ -EA using standard bit mutation with arbitrary mutation rates c/n is $(1 \pm o(1)) \frac{e}{c} n \ln n = (1 \pm o(1)) \frac{1}{p_1} n \ln n$, where p_1 is the probability that the mutation flips a single bit. Our proof of the upper bounds closely follows his, but we need to adapt his potential function to account for the fact that the probabilities p_i may follow any distribution.

For the lower bound we follow the proof strategy from [DDY20]. In particular, we use the same symmetrized ONEMAX potential $X_t = \min_x \min\{\text{OM}(x), n - \text{OM}(x)\}$, where $\text{OM}(x)$ is the number of one-bits in x and the minimum goes over all previously visited search points x . We show that for a wide range of values of X_t , the drift is maximized either by single-bit flips or by $(n - 1)$ -bit flips, and with a parent that achieves the minimum in X_t . This allows us to compute an upper bound on the drift, and to use the variable drift lower bound from [DDY20]. We obtain a lower bound for any function with unique local optimum, but then p_1 needs to be replaced by $p_1 + p_{n-1}$ in (1). This is not an artifact of our proof, since we give examples showing that the dependence on p_{n-1} is real.

Finally, we also show (Section 4.2) that stochastic domination no longer applies when standard bit mutation is replaced by other unary unbiased mutation operators, in the sense that starting closer to the optimum can increase the expected runtime asymptotically. This even holds on ONEMAX. As a consequence, non-elitist algorithms may be faster than elitist algorithms on ONEMAX.

Other Related Work. Apart from black-box complexities, only few things are known in general about the class of unbiased mutation operators. Antipov and Doerr [AD21] investigated the mixing time on plateaus for the $(1 + 1)$ -EA with arbitrary unbiased mutation operator. Lengler [Len19] studied the $(1 + 1)$ -EA, the $(1 + \lambda)$ -EA, the $(\mu + 1)$ -EA and the $(\mu + 1)$ -GA with arbitrary unbiased mutation operators on monotone functions. He found that those algorithms can optimize all monotone functions if the second moment of the number of bit flips is small compared to the first moment, but that they need exponential time on HOTTOPIC functions otherwise. In particular, all heavy-tail distributions lead to exponential runtimes on HOTTOPIC.

2 Preliminaries

We use the following notation. For $a, b \in \mathbb{N}$ with $a \leq b$ we write $[a, b] = \{a, a + 1, \dots, b\}$ and $[b] = [1, b] = \{1, \dots, b\}$. We write a vector $x \in \{0, 1\}^n$ as $x = (x_1, \dots, x_n)$. The ONEMAX value $\text{OM}(x) := \sum_{i=1}^n x_i$ of x is the number of one-bits in x . We write $\vec{0}$ and $\vec{1}$ for the vectors in $\{0, 1\}^n$ with $\text{OM}(\vec{0}) = 0$ and $\text{OM}(\vec{1}) = n$, respectively. *With high probability (w.h.p.)* means with probability $1 - o(1)$ as $n \rightarrow \infty$.

We identify probability distributions $\mathcal{D}$ on $[0, n]$ with sequences $(p_0, p_1, \dots, p_n)$ such that $p_i \geq 0$ for all $i \in [0, n]$ and $\sum_{i \in [0, n]} p_i = 1$, where the probability of obtaining i from $\mathcal{D}$ is p_i . We associate to any such distribution $\mathcal{D}$ the mutation operator $\text{mut}_{\mathcal{D}}$ which draws k from $\mathcal{D}$ and then applies the flip_k operator which flips a uniform random set of exactly k positions. The probability that $\text{mut}_{\mathcal{D}}$ flips the i -th bit equals χ/n , where χ is the expected value of $\mathcal{D}$, as we show in the following lemma.

Lemma 2. *For any $\mathcal{D}$ with mean χ and every $i \in [n]$, the associated mutation operator satisfies*

$$\Pr[i\text{-th bit flips}] = \chi/n. \quad (2)$$

Proof. Let K be the number of flipped bits. Using the law of total probability, the probability that the i -th bit flips is equal to

$$\sum_{k=0}^n \Pr[i\text{-th bit flips} \mid K = k] \cdot p_k = \sum_{k=0}^n \frac{k}{n} \cdot p_k = \frac{\chi}{n}. \quad (3)$$

□

Algorithm 1: The $(1 + 1)$ -EA $_{\mathcal{D}}$ for a fixed distribution $\mathcal{D}$ and maximizing a function $f : \{0, 1\}^n \rightarrow \mathbb{R}$.

```

1 Sample  $x$  from  $\{0, 1\}^n$  uniformly at random;
2 for  $t = 0, 1, 2, 3, \dots$  do
3   Sample  $k \sim \mathcal{D}$ ;
4   Create  $y \leftarrow \text{flip}_k(x)$ ;
5   if  $f(y) \geq f(x)$  then  $x \leftarrow y$ ;
```

For a probability distribution $\mathcal{D}$ on $[0, n]$, we define the $(1 + 1)$ -EA $_{\mathcal{D}}$ as the elitist $(1 + 1)$ algorithm which uses $\text{mut}_{\mathcal{D}}$ as mutation operator, see Algorithm 1. Its *runtime* on a function f is the number of fitness evaluations before it finds a global optimum. Following the discussion in [DDY20, Doe20], the class of elitist $(1+1)$ unary unbiased black-box algorithms with static mutation operators coincides exactly with the collection of all $(1 + 1)$ -EA $_{\mathcal{D}}$ with $\mathcal{D}$ as above.

We call any population based algorithm that generates offspring exclusively using the operator $\text{mut}_{\mathcal{D}}$ a *static unary unbiased algorithm with flip distribution $\mathcal{D}$* . In particular, such an algorithm is not required to use elitist selection, may access any previously generated search point, and is allowed to generate more than one offspring per generation. By using the adjective static, we emphasize that the distribution $\mathcal{D}$ may not change throughout the run of the algorithm.

With *linear functions* we always refer to functions $f : \{0, 1\}^n \rightarrow \mathbb{R}$; $f(x) = \sum_{i=1}^n w_i x_i$ for non-zero weights $w_i \in \mathbb{R}$. By unbiasedness of the $(1 + 1)$ -EA $_{\mathcal{D}}$, we may (and will) assume that the weights are positive and sorted, $0 < w_1 \leq \dots \leq w_n$.

In the following, we briefly recall the mathematical tools needed for our analysis.

2.1 Drift Analysis

As it is the case for Witt's result [Wit13], our upper bound heavily relies on potential function arguments, which are converted into upper bound using the multiplicative drift theorem.

Theorem 3 (Multiplicative Drift Theorem [DJW12]). *Let $S \subset \mathbb{R}$ be a finite set with minimum 1. Let $(X_t)_{t \geq 0}$ be a sequence of random variables over $S \cup \{0\}$. Let $T := \min\{t \geq 0 \mid X_t = 0\}$ be the hitting time of 0. Suppose that there is a real number $\delta > 0$ such that*

$$\mathbb{E}[X_t - X_{t+1} \mid X_t = s] \geq \delta s \quad (4)$$

for all $s \in S$ and all $t \geq 0$ with $\Pr[X_t = s] > 0$. Then, for all $s_0 \in S$ with $\Pr[X_0 = s_0] > 0$,

$$\mathbb{E}[T \mid X_0 = s_0] \leq \frac{\ln(s_0) + 1}{\delta}. \quad (5)$$

Moreover, for all $r > 0$,

$$\Pr\left[T > \frac{\ln(s_0) + r}{\delta}\right] \leq e^{-r}. \quad (6)$$

In the proof of the lower bound we will apply the following lower bound for variable drift [DDY20, Theorem 9].

Theorem 4 (Variable Drift, lower bound). *Let $(X_t)_{t \geq 0}$ be a sequence of non-increasing random variables over $[0, n]$, i.e., it holds $\Pr[X_t \leq X_{t-1}] = 1$ for all $t > 0$, and let $T := \min\{t \geq 0 \mid X^{(t)} = 0\}$ be the hitting time of 0. Suppose that there are two functions $c : [n] \rightarrow [0, n]$ and monotonically increasing $h : [0, n] \rightarrow \mathbb{R}_0^+$, and a constant $0 \leq p < 1$ such that*

1. $X_{t+1} \geq c(X_t)$ with probability at least $1 - p$ for all $t < T$, and
2. $\mathbb{E}[X_t - X_{t+1} \mid X_t] \leq h(X_t)$ for all $t < T$.

Let $\mu : [0, n] \rightarrow [0, n]$ be defined by $\mu(x) := \max\{i \mid c(i) \leq x\}$, and let $g : [0, n] \rightarrow \mathbb{R}_0^+$ be defined by $g(x) := \sum_{i=0}^{x-1} \frac{1}{h(\mu(i))}$. Then

$$\mathbb{E}[T \mid X_0] \geq g(X_0) - \frac{g^2(X_0)p}{1 + g(X_0)p}. \quad (7)$$

2.2 Concentration Bounds

In the proofs for Section 4, we make use of the following additive and multiplicative Chernoff bounds, originally shown by Hoeffding [Hoe63].

Theorem 5 (Additive Chernoff Bound). *Assume that X is a hypergeometrically distributed random variable with parameters N, n, m , or let X be a sum of n independent random variables $X_1, \dots, X_n$, with each taking values in $\{0, 1\}$. Then we have for all $\varepsilon > 0$*

$$\Pr[X \geq \mathbb{E}[X] + \varepsilon] \leq e^{-2\varepsilon^2/n}, \quad \text{and} \quad (8)$$

$$\Pr[X \leq \mathbb{E}[X] - \varepsilon] \leq e^{-2\varepsilon^2/n} \quad (9)$$

Theorem 6 (Multiplicative Chernoff Bound). *Assume that X is a hypergeometrically distributed random variable with parameters N, n, m . Then we have for all $\delta > 0$,*

$$\Pr[X \geq (1 + \delta) \mathbb{E}[X]] \leq \left(\frac{e^\delta}{(1 + \delta)^{1+\delta}} \right)^{\mathbb{E}[X]}. \quad (10)$$

3 Upper Bounds and Tightness Results

We first note the following, simpler version of the upper bound stated in Theorem 1 for ONEMAX, which does not require any assumption on the distribution $\mathcal{D}$. It straightforwardly follows from the standard multiplicative drift theorem [DJW12], applied to the lower bound on the drift obtained by considering only 1-bit flips.

Theorem 7. *Let $\mathcal{D} = (p_0, p_1, \dots, p_n)$ be a probability distribution on $[0, n]$. The runtime of the $(1 + 1)$ -EA $_{\mathcal{D}}$ on ONEMAX is at most*

$$(1 \pm o(1)) \frac{1}{p_1} n \ln n \quad (11)$$

in expectation and with high probability.

Proof. We consider $X_t := n - \text{ONEMAX}(x^{(t)})$. We have

$$\mathbb{E}[X_t - X_{t+1} \mid X_t = s] \geq p_1 \cdot \frac{s}{n}. \quad (12)$$

By Theorem 3, we have

$$\mathbb{E}[T \mid X_0] \leq \frac{\ln n + 1}{p_1/n} = (1 \pm o(1)) \frac{1}{p_1} n \ln n. \quad (13)$$

Taking $r := \ln \ln n$ in Theorem 3 concludes the proof. $\square$

The key ingredient for generalizing the bound from ONEMAX to all linear functions as in Theorem 1 is the following theorem, which generalizes [Wit13, Theorem 4.1] to the $(1+1)$ -EA $_{\mathcal{D}}$ with (almost) arbitrary $\mathcal{D}$. Our proof follows [Wit13], with the following differences: First, we noted a mistake in the proof of the upper bound in [Wit13]. Equation (4.2) there does not hold for the events A_i as defined in [Wit13]. We thank Carsten Witt for providing the following fix upon our inquiry (personal communication): By conditioning the events A_i on the event that the offspring is accepted, equation (4.2) holds as in that case, the expectation is zero if none of the A_i occur. Furthermore, the inequality (4.3) in [Wit13] still holds, which can be shown by applying Bayes' theorem and linearity of expectation.

Apart from this issue, the biggest challenge was to adapt the potential function used in [Wit13], since we need to deal with arbitrary unbiased mutation operators. In particular, our potential involves the quantities χ and p_1 . With the modified potential, we can show the following generalization of [Wit13, Theorem 4.1].

Theorem 8. *Let $\mathcal{D} = (p_0, p_1, \dots, p_n)$ be a probability distribution on $[0, n]$ with expectation χ and with $p_1 > 0$. Then the runtime of the $(1+1)$ -EA $_{\mathcal{D}}$ on any linear function on n variables is at most*

$$b(r) := \frac{n}{p_1} \cdot \frac{\alpha}{\alpha - 1} \cdot \left(\frac{\alpha n \chi^3}{(n-1)p_1^2} + \ln \left(\frac{(n-1)p_1^2}{\chi^3} \right) + r \right) \quad (14)$$

with probability at least $1 - e^{-r}$ for any $r > 0$, and it is at most $b(1)$ in expectation, where $\alpha > 1$ can be chosen arbitrarily.

To ease the comparison of our proof of Theorem 8 and Theorem 4.1 in [Wit13], we keep the same notation. For the parts of the proof in [Wit13] that transfer directly to our case, we will simply cite them.

Proof of Theorem 8. Let $f : \{0, 1\}^n \rightarrow \mathbb{R}$, $f(x) = w_1 x_1 + \dots + w_n x_n$, with $0 < w_1 \leq \dots \leq w_n$. The proof works by applying the multiplicative drift theorem to a carefully chosen potential. To this end, following [Wit13], we define a new (linear) function g , and consider the stochastic process $X_t = g(x^{(t)})$, where $x^{(t)}$ is the current search point of the $(1+1)$ -EA $_{\mathcal{D}}$ at time t . The weights g_i of the function g are as follows. For all $1 \leq i \leq n$, we let

$$\gamma_i := \left(1 + \frac{\alpha \chi^3}{(n-1)p_1^2} \right)^{i-1}, \quad (15)$$

put $g_1 := \gamma_1 = 1$, and for $2 \leq i \leq n$ we set

$$g_i := \min \left\{ \gamma_i, g_{i-1} \frac{w_i}{w_{i-1}} \right\} \geq 1. \quad (16)$$

We note that the g_i are non-decreasing with respect to i , and define $g(x) := g_1 x_1 + \dots + g_n x_n$ and $X_t := g(x^{(t)})$. Then $X_t = 0$ if and only if f has been optimised. Let $\Delta_t := X_t - X_{t+1}$. First, we will show that

$$\mathbb{E}[\Delta_t \mid X_t = s] \geq s \cdot \frac{p_1}{n} \cdot \frac{\alpha - 1}{\alpha}. \quad (17)$$

In the following, we recall some notation from [Wit13]. Fix $s \in [0, n]$ and a search point $x^{(t)}$ with $g(x^{(t)}) = s$. From now on, we implicitly assume that $X_t = s$. Let $I := \{i \mid x_i^{(t)} = 1\}$ be the index set of the one-bits in $x^{(t)}$ and $Z := [n] \setminus I$ be the zero-bits. Let x' denote the random offspring generated by the $(1+1)$ -EA $_{\mathcal{D}}$ from $x^{(t)}$, and let $x^{(t+1)}$ be the next search point after selection. We denote by $I^* := \{i \in I \mid x'_i = 0\}$ the index set of one-bits that are flipped, and by $Z^* := \{i \in Z \mid x'_i = 1\}$ the zero-bits. Let $k(i) := \max\{j \leq i \mid g_j = \gamma_j\}$ for all $i \in I$. Note that $k(i) \geq 1$. Set $L(i) := [k(i), n] \cap Z$ and $R(i) := [1, k(i) - 1] \cap Z$.

For $i \in I$, we define the events A_i as

1. i is the leftmost flipping one-bit, and
2. $\sum_{j \in I^*} w_j - \sum_{j \in Z^*} w_j \geq 0$, i.e. the offspring is accepted.

Note that the A_i are mutually disjoint. Furthermore, if none of the A_i occur we have $\Delta_t = 0$, as then the offspring is either rejected or equal to the parent. Let

$$\Delta_L(i) := \sum_{j \in I^*} g_j - \sum_{j \in Z^* \cap L(i)} g_j, \text{ and} \quad (18)$$

$$\Delta_R(i) := - \sum_{j \in Z^* \cap R(i)} g_j. \quad (19)$$

Conditioning on A_i , we have that $\Delta_t = \Delta_L(i) + \Delta_R(i)$, as in that case, the offspring is accepted. By the law of total expectation,

$$\begin{aligned}\mathbb{E}[\Delta_t] &= \sum_{i \in I} \mathbb{E}[\Delta_t \mid A_i] \cdot \Pr[A_i] + \underbrace{\mathbb{E}\left[\Delta_t \mid \bigcup_{i \in I} A_i\right]}_{=0} \cdot \Pr\left[\bigcup_{i \in I} A_i\right] \\ &= \sum_{i \in I} \mathbb{E}[\Delta_L(i) + \Delta_R(i) \mid A_i] \cdot \Pr[A_i] \\ &= \sum_{i \in I} \mathbb{E}[\Delta_L(i) \mid A_i] \cdot \Pr[A_i] + \mathbb{E}[\Delta_R(i) \mid A_i] \cdot \Pr[A_i],\end{aligned}\tag{20}$$

where we used linearity of conditional expectation for the last step.

In the following, we want to estimate the terms appearing in the above sum. First, we turn our attention to $\mathbb{E}[\Delta_L(i) \mid A_i] \cdot \Pr[A_i]$. It follows from the same argument as [Wit13] uses to show the non-negativity of $\Delta_L(i)$ that $\mathbb{E}[\Delta_L(i) \mid A_i] \geq 0$. As in [Wit13], for the event $S_i := \{Z^* \cap L(i) = \emptyset\}$, we get

$$\mathbb{E}[\Delta_L(i) \mid A_i] \Pr[A_i] \geq g_i \cdot \Pr[A_i \cap S_i].\tag{21}$$

To lower bound $\Pr[A_i \cap S_i]$, we note that $A_i \cap S_i$ occurs if we flip exactly the i -th bit. The probability of flipping exactly the i -th bit, which is p_1/n , is thus a lower bound for $\Pr[A_i \cap S_i]$, yielding

$$\mathbb{E}[\Delta_L(i) \mid A_i] \cdot \Pr[A_i] \geq g_i \cdot \frac{p_1}{n}.\tag{22}$$

In the following lemma, we estimate $\mathbb{E}[\Delta_R(i) \mid A_i]$.

Lemma 9. *We have*

$$\mathbb{E}[\Delta_R(i) \mid A_i] \geq -\frac{\chi^2}{(n-1)p_1} \sum_{j \in R(i)} g_j.\tag{23}$$

Proof. It holds

$$\begin{aligned}\mathbb{E}[\Delta_R(i) \mid A_i] &= \mathbb{E}\left[-\sum_{j \in Z^* \cap R(i)} g_j \mid A_i\right] = \mathbb{E}\left[-\sum_{j \in R(i)} \mathbb{1}_{\{j \in Z^*\}} g_j \mid A_i\right] \\ &= -\sum_{j \in R(i)} g_j \Pr[\{j \in Z^*\} \mid A_i],\end{aligned}\tag{24}$$

where we used linearity and the definition of conditional expectation. Since for all $i \in I$, $\Pr[A_i] > 0$ (it is possible to flip just the i -th bit, as $p_1 > 0$), we can apply Bayes' theorem to the conditional probability above, yielding

$$\Pr[\{j \in Z^*\} \mid A_i] = \Pr[\{j \in Z^*\}] \cdot \frac{\Pr[A_i \mid \{j \in Z^*\}]}{\Pr[A_i]}.\tag{25}$$

By Lemma 2, $\Pr[\{j \in Z^*\}] = \chi/n$. We lower bound the denominator by p_1/n . The numerator is at most

$$\Pr[i\text{-th bit flips} \mid \{j \in Z^*\}],\tag{26}$$

as it is necessary to flip the i -th bit for A_i to occur. We calculate using the law of total probability

$$\Pr[i\text{-th bit flips} \mid \{j \in Z^*\}] = \Pr[i\text{-th bit flips} \mid j\text{-th bit flips}]$$

$$\begin{aligned}
&= \sum_{k=2}^n \Pr[i\text{-th bit flips} \mid j\text{-th bit flips} \cap \{k \text{ bits flip in total}\}] \cdot p_k \\
&= \sum_{k=2}^n \frac{k-1}{n-1} \cdot p_k = \frac{1}{n-1} \left(\sum_{k=1}^n (k-1)p_k \right) \leq \frac{\chi}{n-1}.
\end{aligned} \tag{27}$$

Plugging the bounds obtained above into (25) and (24) completes the proof of Lemma 9. $\square$

We now continue with the proof of Theorem 8. By Lemma 2, we have $\Pr[A_i] \leq \Pr[i\text{-th bit flips}] = \chi/n$, as the i -th bit needs to flip for A_i to occur. Plugging this, Lemma 9, and (22) into (20), we get

$$\begin{aligned}
\mathbb{E}[\Delta_t] &\geq \sum_{i \in I} \left(g_i \cdot \frac{p_1}{n} - \frac{\chi}{n} \cdot \frac{\chi^2}{(n-1)p_1} \sum_{j \in R(i)} g_j \right) \\
&\geq \sum_{i \in I} \left(\frac{p_1}{n} \cdot \frac{g_i}{g_{k(i)}} \cdot \gamma_{k(i)} - \frac{\chi^3}{n(n-1)p_1} \sum_{j=1}^{k(i)-1} \gamma_j \right).
\end{aligned} \tag{28}$$

It was shown on page 304 of [Wit13] that the i -th summand in (28) is at least

$$\frac{\alpha-1}{\alpha} \cdot \frac{p_1}{n} \cdot g_i. \tag{29}$$

Plugging this back into (28) gives us

$$\mathbb{E}[\Delta_t] \geq \sum_{i \in I} \frac{\alpha-1}{\alpha} \cdot \frac{p_1}{n} \cdot g_i = \frac{\alpha-1}{\alpha} \cdot \frac{p_1}{n} \cdot g(x^{(t)}), \tag{30}$$

which shows (17).

Finally, we apply the multiplicative drift theorem (Theorem 3) to finish the proof. To this end, we compute

$$X_0 \leq \sum_{i=1}^n g_i \leq \sum_{i=1}^n \gamma_i = \frac{\left(1 + \frac{\alpha\chi^3}{(n-1)p_1^2}\right)^n - 1}{\frac{\alpha\chi^3}{(n-1)p_1^2}} \leq \frac{e^{n \cdot \frac{\alpha\chi^3}{(n-1)p_1^2}}}{\frac{\alpha\chi^3}{(n-1)p_1^2}}. \tag{31}$$

Hence,

$$\ln X_0 \leq n \cdot \frac{\alpha\chi^3}{(n-1)p_1^2} + \ln \left(\frac{(n-1)p_1^2}{\chi^3} \right), \tag{32}$$

as $-\ln \alpha < 0$ (because $\alpha > 1$). We apply Theorem 3 with $\delta = ((\alpha-1)/\alpha) \cdot (p_1/n)$, which concludes the proof of Theorem 8. $\square$

From Theorem 8 we obtain the following upper bound, which relaxes the conditions on p_1 and χ in Theorem 1 and shows that the bound holds not only in expectation but also w.h.p.

Corollary 10. *Let $\mathcal{D} = (p_0, p_1, \dots, p_n)$ be a probability distribution on $[0, n]$ with expectation χ . Assume that $p_1 > 0$ and $\chi^3 p_1^{-2} (1 - p_0)^{-1} = o(\ln n / \ln \ln n)$. Then the runtime of the $(1 + 1)$ -EA $_{\mathcal{D}}$ on any linear function is at most*

$$(1 + o(1)) \frac{1}{p_1} \cdot n \ln n \tag{33}$$

in expectation and with high probability.

Note that since $1 - p_0 \geq p_1$, we could replace the requirement $\chi^3 p_1^{-2} (1 - p_0)^{-1} = o(\ln n / \ln \ln n)$ by the stronger requirement $\chi^3 / p_1^3 = o(\ln n / \ln \ln n)$. In particular, this is trivially satisfied if $p_1 = \Theta(1)$ and $\chi = O(1)$, as required in Theorem 1.

Proof. We first treat the case $p_0 = 0$, which implies $\chi \geq 1$. Let $\alpha := \ln \ln n$. As in [Wit13], $\alpha/(\alpha - 1) = 1 + O(1/\ln \ln n)$, and $\alpha^2/(\alpha - 1) = O(\ln \ln n)$. Moreover, we may bound $\ln((n - 1)p_1^2/\chi^3) \leq \ln n$, as $p_1 \leq 1$ and $\chi \geq 1$. Thus $b(r)$ in Theorem 8 is at most

$$\frac{n}{p_1} (o(\ln n) + (1 + o(1)) (\ln n + o(\ln n) + r)) \quad (34)$$

Taking $r := 1$, we get the claimed expected runtime, and with $r := \ln \ln n$, it follows that the bound holds w.h.p.

Now we turn to $p_0 > 0$. In this case, we define an auxiliary distribution $\mathcal{D}' = (p'_0, p'_1, \dots, p'_n)$ by $p'_0 := 0$ and $p'_i := p_i/(1 - p_0)$ for $i \geq 1$. In other words, $\mathcal{D}'$ is the same as the distribution $\mathcal{D}$ conditioned on not drawing 0. It has expectation $\chi' := \chi/(1 - p_0)$. Therefore,

$$\frac{(\chi')^3}{(p'_1)^2(1 - p'_0)} = \frac{\chi^3}{p_1^2(1 - p_0)} = o\left(\frac{\ln n}{\ln \ln n}\right). \quad (35)$$

Hence, $\mathcal{D}'$ is covered by the case that we have already treated. Thus, the runtime of the $(1 + 1) - EA_{\mathcal{D}'}$ is at most

$$(1 + o(1)) \frac{1}{p'_1} \cdot n \ln n = (1 + o(1)) \frac{1 - p_0}{p_1} \cdot n \ln n, \quad (36)$$

in expectation and with high probability.

Note that no-bit flips are just idle steps of the $(1 + 1) - EA_{\mathcal{D}}$, therefore the $(1 + 1) - EA_{\mathcal{D}}$ and the $(1 + 1) - EA_{\mathcal{D}'}$ follow exactly the same trajectory through the search space, except that the $(1 + 1) - EA_{\mathcal{D}}$ performs idle steps with probability p_0 . Note that the expected time until a non-idle step is $1/(1 - p_0)$. Hence, if the $(1 + 1) - EA_{\mathcal{D}'}$ finds the optimum in T steps, then the $(1 + 1) - EA_{\mathcal{D}}$ needs $\frac{1}{(1 - p_0)}T$ steps in expectation, and $(1 + o(1))\frac{1}{(1 - p_0)}T$ steps with high probability. Together with (36), this implies the claim. $\square$

Under some less restrictive conditions, we can give a polynomial upper bound on the expected runtime. **Corollary 11.** *Let $\mathcal{D} = (p_0, p_1, \dots, p_n)$ be a probability distribution on $[0, n]$ with expectation χ . The runtime of the $(1 + 1) - EA_{\mathcal{D}}$ on any linear function is $O(n\chi^3/p_1^3 + n \ln n/p_1)$ in expectation and with high probability. In particular, this expression is $O(n^4/p_1^3)$, and thus polynomial in n if $1/p_1$ is polynomial in n .*

Proof. We take $\alpha := 2$, and apply Theorem 8. Noting that $\chi \geq p_1$ and thus $(n - 1)p_1^2/\chi^3 \leq n/p_1$, by (14) we can bound the runtime of the $(1 + 1) - EA_{\mathcal{D}}$ on any linear function by

$$\frac{2n}{p_1} \left(\frac{2n\chi^3}{(n - 1)p_1^2} + \ln\left(\frac{n}{p_1}\right) + r \right) = O\left(\frac{n\chi^3}{p_1^3} + \frac{n \ln n}{p_1} + r\right), \quad (37)$$

where the bound holds in expectation for $r = 1$, and with high probability for e.g. $r = \ln n$. This proves the first statement. The second statement holds since $\chi \leq n$. $\square$

3.1 Tightness

We now discuss that some requirements on p_1 and χ are necessary.

Requirement on p_1 . We start with a proposition saying that the leading constant can change if $p_1 = n^{-c}$ for any $c > 0$. In fact, this is already the case for ONEMAX, as the following example shows.

Proposition 12. *Let $0 < c < 1$ be constant. Consider the $(1 + 1) - EA_{\mathcal{D}}$ with distribution $\mathcal{D} = (p_0, p_1, \dots, p_n)$ defined by $p_1 = n^{-c}$ and $p_2 = 1 - p_1$. Then there is $\varepsilon > 0$ such that for sufficiently large n the expected runtime on ONEMAX is at most*

$$(1 - \varepsilon) \cdot \frac{1}{p_1} \cdot n \ln n.$$

Proof. We consider two phases of the algorithm: one for reducing the distance from the optimum from n to $d_0 := n^{1-c/2}$, and the other for reducing it from d_0 to 0. Let T_1 and T_2 be the respective time spent in those two phases.

To bound T_1 , assume the current distance is $d > d_0$. Then the probability that both bits of a 2-bit flip are zero-bits is $\binom{d}{2}/\binom{n}{2} \geq d_0^2/n^2 = n^{-c}$. We need at most $(n - d_0)/2$ such flips to leave the first phase, and thus

$$\mathbb{E}[T_1] \leq \frac{n - d_0}{2} \cdot n^c \leq \frac{n^{1+c}}{2} = o\left(\frac{1}{p_1} n \ln n\right) \quad (38)$$

rounds. Hence, the first phase is asymptotically faster than the claimed runtime bound.

For the second phase, if X_t is the distance from the optimum in round t , then every single-bit flip has a chance of X_t/n to reduce X_t by one. Since single-bit flips occur with probability p_1 , we have

$$\mathbb{E}[X_t - X_{t+1} \mid X_t = d] \geq p_1 \cdot \frac{d}{n}. \quad (39)$$

The second phase starts with $X_0 \leq d_0$, so by the multiplicative drift theorem the duration of the second phase is at most

$$\mathbb{E}[T_2] \leq \frac{1 + \ln(d_0)}{p_1/n} = \frac{n + (1 - c/2)n \ln n}{p_1}, \quad (40)$$

where we used $\ln(d_0) = (1 - c/2) \ln n$. Taken together, we obtain

$$\mathbb{E}[T_1 + T_2] \leq (1 + o(1)) \cdot \frac{1 - c/2}{p_1} \cdot n \ln n, \quad (41)$$

and the proposition follows with $\varepsilon := c/3$. $\square$

The point of Proposition 12 is that we have a strictly smaller leading constant than in Corollary 10 and in Theorem 1. The reason for this effect is that if $p_1 = n^{-\Omega(1)}$, for Hamming distances in the range $[n^{1-c/2}, n]$ from the optimum, two-bit flips are more effective than single-bit flips. This range is thus traversed more quickly. With single-bit flips, the algorithm would need time $\Omega(n \ln n / p_1)$ to traverse this region, but it can be traversed in time $o(n \ln n / p_1)$ by two-bit flips. Hence, the time spent in this phase becomes negligible. Even though this region is still far away from the optimum, it consumes a constant fraction of the total expected runtime if the algorithm is restricted to single-bit flips. Hence, the speed-up from two-bit flips eliminates this constant fraction from the total expected runtime, and thus reduces the leading constant of the total expected runtime. This shows that the lower bound in Theorem 1 can not hold if $p_1 = n^{-\Omega(1)}$. On the other hand, we will show in Theorem 23 below that it does hold for all $p_1 = n^{-o(1)}$, which is tight by the above discussion.

Requirement on χ . Other than for p_1 , we could not derive a statement about the runtime, but the following proposition shows that, close to the optimum, the behavior of the algorithm changes substantially if χ is large. Recall that from any parent at distance one from the optimum, we have a probability of $p_1 \cdot 1/n$ to create the optimum as offspring. Hence, one would naively expect to wait at most for n/p_1 rounds in expectation to find the optimum. However, this is wrong for large values of χ , as the following proposition shows.

Proposition 13. *Let x be a search point at Hamming distance one from the optimum $\vec{1}$. Let $\mathcal{D}$ be any probability distribution on $[0, n]$ with mean χ . For a linear function f , let $T_{\mathcal{D}}^x(f)$ be the number of iterations until the $(1 + 1)$ -EA $_{\mathcal{D}}$ with starting position x finds the optimum.*

- (a) *There is a linear function f depending on x such that $\mathbb{E}[T_{\mathcal{D}}^x(f)] = \Omega(n \ln \chi)$.*
- (b) *If $\chi = \omega(1)$ then there is a linear function f depending on x such that $T_{\mathcal{D}}^x(f) = \omega(n)$ with high probability.*
- (c) *If $\chi = O(1)$ and $p_1 = \Theta(1)$ then $\mathbb{E}[T_{\mathcal{D}}^x(f)] = O(n)$ for every linear function f .*

Proof. (a) and (b). It is clear that the number of iterations is at least $\Omega(n)$, so we may assume $\chi \geq 4$. Since x has Hamming distance one from $\vec{1}$, it differs in exactly one position from $\vec{1}$. We may assume

that this is the first position. Then we define f via $f(x) := n \cdot x_1 + \sum_{i=2}^n x_i$, i.e., we give weight n to the first position and weight 1 to all other positions. When in x , the algorithm will accept any offspring that flips the first position. Let us call R the number of bits that are flipped in this mutation. Then we may compute the distribution of R via Bayes formula as

$$\Pr[R = r] = \frac{\Pr[R = r \text{ and pos. 1 flipped}]}{\sum_{s \in [n]} \Pr[R = s \text{ and pos. 1 flipped}]}. \quad (42)$$

Note that $\Pr[R = s \text{ and pos. 1 flipped}] = \Pr[R = s] \cdot \Pr[\text{pos 1 flipped} \mid R = s] = p_s \cdot s/n$, where the conditional probability is s/n since the mutation operator is unbiased. Hence, (42) simplifies to

$$\Pr[R = r] = \frac{p_r \cdot r/n}{\sum_{s \in [n]} p_s \cdot s/n} = \frac{p_r \cdot r}{\chi}. \quad (43)$$

In particular, this implies for every $\gamma \leq 1$,

$$\Pr[R \leq \gamma\chi] = \sum_{r=1}^{\lfloor \gamma\chi \rfloor} \frac{p_r \cdot r}{\chi} \leq \frac{\gamma\chi \sum_{r=1}^{\lfloor \gamma\chi \rfloor} p_r}{\chi} \leq \gamma.$$

Hence, with probability at least $1 - \gamma$, the $(1 + 1)\text{-EA}_{\mathcal{D}}$ proceeds from x to a search point in Hamming distance at least $\gamma\chi - 1$ from $\vec{1}$.

For claim (a) we set $\gamma := 1/2$ and obtain a search point in Hamming distance at least $\chi/2 - 1 \geq \chi/4$. Once this search point is reached, the algorithm does not accept any mutation which flips position 1 again, since this would decrease the fitness. Hence, the algorithm simply has to solve ONEMAX on the remaining $n - 1$ bits. In expectation, this takes time $\Omega(n \ln \chi)$ since the unbiased black-box complexity for solving ONEMAX from a starting point in distance $d > 1$ from the optimum is $\Omega(n \ln d)$ (implicit in [LW12]). Since this case happens with probability at least $1 - \gamma = 1/2$, we obtain

$$\mathbb{E}[T_{\mathcal{D}}^x(f)] \geq \frac{1}{2} \cdot \Omega(n \ln \chi) = \Omega(n \ln \chi).$$

For claim (b), we choose $\gamma := \chi^{-1/2}$. Then $1 - \gamma = 1 - o(1)$, so with high probability the $(1 + 1)\text{-EA}_{\mathcal{D}}$ proceeds from x to a search point in distance at least $d := \gamma\chi - 1 = \chi^{1/2} - 1 = \omega(1)$ from $\vec{1}$. As for part (a), from this point onwards the algorithm needs to solve ONEMAX on $n - 1$ bits. Thus the algorithm needs to traverse the interval from $d' := \min\{d, n^{1/3}\}$ to the optimum. We show that w.h.p. this takes time at least t_0 for some $t_0 = \Omega(n \ln d') = \omega(n)$. It can be shown that the probability to find an improvement with r -bit flips for any $r \geq 2$ is $O((d'/n))^2 \leq n^{-4/3}$. Hence, in time t_0 the expected number of such improvements is $O(t_0 n^{-4/3}) = o(1)$, and by Markov's inequality no r -bit flip finds an improvement for $r \geq 2$. Hence, we may pessimistically assume that the algorithm only uses single-bit flips, i.e., that it is random local search (RLS). By [Wit14, Theorem 1], w.h.p. RLS needs time $\Omega(n \ln d')$ to find the optimum, which concludes the proof.

Claim (c) follows directly from the proof of Theorem 8, using the parameter $\alpha = 2$. There it was shown that with the potential $g(x) = \sum_{i=1}^n g_i x_i$, the drift towards the optimum is at least $\frac{p_1}{2n} \cdot g(x)$ by (17). By design, the minimal positive potential is 1. Since we start in distance one from the optimum, the initial potential is at most

$$g_{\text{init}} \leq \max\{g_i : i \in [n]\} = g_n \leq \gamma_n, \quad (44)$$

where

$$\gamma_n = \left(1 + \frac{2\chi^3}{(n-1)p_1^2}\right)^{n-1} \leq \exp\left(\frac{2\chi^3}{p_1^2}\right) = O(1) \quad (45)$$

by (15) and (16). Hence, the expected runtime is at most

$$\mathbb{E}[T] \leq \frac{\ln(g_{\text{init}}) + 1}{p_1/(2n)} = \frac{O(n)}{p_1} = O(n) \quad (46)$$

by the multiplicative drift theorem. $\square$

We did not aim for tightness in Proposition 13, but rather want to demonstrate the different regimes. In particular, consider a distribution $\mathcal{D}$ with $p_1 = \Theta(1)$ and with mean χ . If $\chi = O(1)$, then $\mathbb{E}[T_{\mathcal{D}}^x(f)] = O(n)$ by (c), but for $\chi = \omega(1)$ we have $\mathbb{E}[T_{\mathcal{D}}^x(f)] = \omega(n)$ by (a). This shows that the size of χ is truly relevant for the runtime, at least if the algorithm starts in an adversarial point. Moreover, (b) shows that the high expectation in the case $\chi = \omega(1)$ is not just due to low-probability events, but that it comes from typical runs.

The most interesting and common unbiased mutation operators, except for standard-bit mutation, are mutation operators where the number of bit flips has a *heavy tail*. Usually a *power law* is used, i.e., the probability to flip k bits scales like $k^{-\alpha}$ for some constant $\alpha > 1$. There are two different regimes for the parameter α . For $\alpha > 2$, the expected number of bit flips satisfies $\chi = O(1)$. For $\alpha \in (1, 2)$, the expected number of bit flips is unbounded and large, $\chi = n^{\Omega(1)}$.² In either case, such power-law distributions satisfy $p_1 = \Theta(1)$. Notably, our main Theorem 1 applies to power-law distributions with $\alpha > 2$, but not to power-law distributions with $\alpha \in (1, 2)$. We believe that this reflects a real difference between those two cases. As indication, note that in the situation of Proposition 13, we have $\mathbb{E}[T_{\mathcal{D}}^x(f)] = O(n)$ for $\alpha > 2$, but $\mathbb{E}[T_{\mathcal{D}}^x(f)] = \Omega(n \ln n)$ for $\alpha \in (1, 2)$. This does not rule out that Theorem 1 still might be true for $\alpha \in (1, 2)$ due to the random starting point, but it suggests that trajectories of the algorithm can be substantially different.

4 Lower Bound

The following theorem is the main result shown in this section.

Theorem 14. *The expected runtime of any static unary unbiased algorithm with flip distribution $\mathcal{D} = (p_0, p_1, \dots, p_n)$ satisfying $p_1 + p_{n-1} = n^{-o(1)}$ on any function $f : \{0, 1\}^n \rightarrow \mathbb{R}$ with unique global optimum is at least*

$$(1 - o(1)) \frac{1}{p_1 + p_{n-1}} n \ln n. \quad (47)$$

Note that most common mutation operators satisfy $p_{n-1} = o(p_1)$, in which case (47) simplifies to $(1 - o(1)) \frac{1}{p_1} n \ln n$. We remark that a coarser lower bound of $\Omega(n \ln n)$ follows from [LW12] and [DDY20]. However, in contrast to [LW12], we are interested in understanding the leading constant, and in contrast to [DDY20], we are interested in the expected runtime for *static* unary unbiased distributions. A common technique to prove lower bounds that apply to any function from some problem collection is to bound the expected runtime of the algorithm on ONEMAX and to show that ONEMAX is the “easiest” among all functions from the collection, in the sense that the expected runtime of the algorithm optimizing a given function from the set cannot be smaller than its expected runtime on ONEMAX. In many cases, e.g., when considering the $(1+1)$ -EA with standard bit mutation, ONEMAX can even be shown to be the easiest among all functions with unique global optimum; as was first shown in [DJW12] for mutation rate $p = 1/n$ and then in [Wit13] more generally for all (static or dynamic) mutation rates $p \leq 1/2$. However, in our situation it is not true that ONEMAX is the easiest function, as we will discuss in Section 4.2.

Our proof for Theorem 14 follows the strategy used in [DDY20]. In particular, we apply their lower bound theorem for variable drift [DDY20, Theorem 9] in the same way. We quote their Lemma 13 directly, and the proof of our Theorem 21 below differs from the proof of their Theorem 14 only in the calculations and bounds used. The key difference between their proof and ours is that we use a different function h to bound the expected change in the potential. Most of the work goes into showing that this function h is indeed applicable, and providing an upper bound on its values that allows us to translate the result of Theorem 21 into the asymptotic formulation of Theorem 14.

²For $\alpha = 2$ the expected number of bit flips is also unbounded, but grows only as $\chi = O(\ln n)$. We will neglect this case here.

To implement the proof strategy of [DDY20], we use the same potential function to measure the progress of the optimization process. That is, we denote by $(x^{(0)}, x^{(1)}, \dots, x^{(t)})$ the sequence of the first $t + 1$ search points evaluated by the algorithm and we denote by $v_t \in \{x^{(0)}, \dots, x^{(t)}\}$ the parent chosen by the algorithm in iteration t . We define the potential at time t as

$$X_t := \min_{0 \leq i \leq t} d(x^{(i)}), \quad (48)$$

where d is the distance function

$$d(x) := \min\{n - \text{OM}(x), \text{OM}(x)\}. \quad (49)$$

The reason for considering the symmetric distance to the optimum and its complement is that an optimal unary unbiased black-box algorithm may first reach $\vec{0}$, and then flip all bits at once. Furthermore, as we show in Lemma 22, it is possible to make progress towards the optimum in a way that can be measured in terms of d , while the Hamming-distance to the optimum increases.

Note that the sequence $(X_t)_{t \geq 0}$ is non-increasing, so we may apply the variable drift lower bound from [DDY20, Theorem 9] to it.

Next, we define the function $\tilde{h}$, which gives the precise drift in the case where the algorithm uses a bitstring at distance X_t for generating the offspring in round t .

Definition 1. We define $\tilde{h} : [0, n] \rightarrow \mathbb{R}_{\geq 0}$ as

$$\tilde{h}(d) = \sum_{r=1}^{n-1} (p_r + p_{n-r}) B(n, d, r), \quad (50)$$

where

$$B(n, d, r) = \sum_{i=\max\{\lceil r/2 \rceil, r+d-n\}}^{\min\{d, r\}} (2i - r) \frac{\binom{d}{i} \binom{n-d}{r-i}}{\binom{n}{r}} \quad (51)$$

is the drift conditioned on flipping r bits.

The expression $B(n, d, r)$ was already given in [DDY20] as the exact fitness drift with respect to ONEMAX when flipping r bits.

Lemma 15. *If a static unary unbiased algorithm with flip distribution $\mathcal{D}$ chooses a bitstring v_t with potential X_t for mutation in step t , then the drift is given by $\tilde{h}(X_t)$, i.e.,*

$$\mathbb{E}[X_t - X_{t+1} \mid \{X_t = d\} \wedge \{d(v_t) = d\}] = \tilde{h}(d). \quad (52)$$

The proof relies on the fact that flipping $n - r$ bits is the same as first flipping n bits and then flipping r bits to adapt the computation of $B(n, d, r)$ given in [DDY20].

Proof. We assume without loss of generality that $d = \text{OM}(v_t)$ (by the symmetry of d , the proof in the case where $d = n - \text{OM}(x)$ is the same after switching the roles of zero-bits and one-bits). The algorithm makes progress if

1. $\text{OM}(x^{(t+1)}) < d$, or
2. $n - \text{OM}(x^{(t+1)}) < d$.

Assume that the algorithm flips r bits in the t -th round. For the first case, let i be the number of bits that flip from 1 to 0, i.e., the number of bits that make progress towards $\vec{0}$. Then $r - i$ bits flip from 0 to 1. The probability of this event is

$$\Pr[\{\text{flip } i \text{ one-bits}\} \mid \{\text{flip } r \text{ bits}\}] = \frac{\binom{d}{i} \binom{n-d}{r-i}}{\binom{n}{r}}. \quad (53)$$

In this event, we make $i - (r - i) = 2i - r$ progress towards $\vec{0}$. This is positive if and only if $i > r/2$. Considering the conditions that $i \leq d$ and $r - i \leq n - d$, we get from the law of total expectation

(conditioning on the number r of bits flipped by the algorithm) for the first case

$$\sum_{r=1}^{n-1} p_r B(n, d, r). \quad (54)$$

Note that flipping $r = 0$ or $r = n$ bits never yields progress with respect to d .

To deal with the second case, we observe that flipping $n - r$ bits is equivalent to first flipping all n bits, and then flipping r bits. Similarly, flipping r bits is the same as first flipping n bits, and then flipping $n - r$ bits. The progress that the algorithm makes towards $\vec{0}$ after first flipping all n bits is the same as the progress the algorithm makes towards $\vec{1}$. Overall, we see that the progress towards $\vec{1}$ of flipping r bits is the same as the progress towards $\vec{0}$ of flipping $n - r$ bits. By the same argument as for the first case, the latter is given by $B(n, d, n - r)$. Hence, the second case contributes

$$\sum_{r=1}^{n-1} p_r B(n, d, n - r). \quad (55)$$

to the drift. Adding up (54) and (55) yields (50), as required. $\square$

Now, we are ready to define the bound h on the drift that we use in our application of the variable drift theorem [DDY20, Theorem 9].

Definition 2. Let $h : [0, n] \rightarrow \mathbb{R}_{\geq 0}$,

$$h(d) = \begin{cases} \tilde{h}(d), & \text{for } d \leq d_0 \\ n, & \text{for } d > d_0, \end{cases} \quad (56)$$

where

$$d_0 := \lfloor (p_1 + p_{n-1})n / \ln^2 n \rfloor. \quad (57)$$

The following statement is adapted from Lemma 21 in [DDY20].

Lemma 16. *There is an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, $d \leq d_0$, and $r \geq r_0 = 12$, it holds*

$$B(n, d, r) < (d/n)^2. \quad (58)$$

Proof. Note that the probability in equation (53) is the same as the probability that a hypergeometric random variable with parameters n, r, d is equal to i . Let X be such a random variable. We have $\mathbb{E}[X] = \frac{dr}{n} \leq (p_1 + p_{n-1}) \frac{r}{\ln^2 n} \leq \frac{r}{\ln^2 n}$.

Let $n_0 = e^{20}$. Then $d/n \leq d_0/n \leq 1/\ln^2 n \leq 1/400$. Making use of the above observation, the fact that $(2i - r) \leq r$ for all $i \leq r$, and then applying Theorem 6, we have

$$\begin{aligned} B(n, d, r) &\leq r \Pr[X \geq r/2] = r \Pr \left[X \geq \left(1 + \left(\frac{n}{2d} - 1 \right) \right) \cdot \mathbb{E}[X] \right] \\ &\leq r \left(\frac{\exp \left(\frac{n}{2d} - 1 \right)}{\left(\frac{n}{2d} \right)^{\frac{n}{2d}}} \right)^{\frac{dr}{n}} \leq r \frac{e^{r/2}}{\left(\frac{n}{2d} \right)^{r/2}} \leq r \left(\frac{2ed}{n} \right)^{r/2} \\ &= 4e^2 r \left(\frac{2ed}{n} \right)^{r/2-2} \cdot \left(\frac{d}{n} \right)^2 < 4e^2 r \left(\frac{e}{200} \right)^{r/2-2} \cdot \left(\frac{d}{n} \right)^2 \\ &< r (2^{-4})^{r/2-3} \cdot (d/n)^2 = \frac{2^{12}r}{2^{2r}} \cdot (d/n)^2 \leq (d/n)^2. \end{aligned} \quad (59)$$

For the last inequality, we used the facts that $r \geq r_0 = 12$ and $r \leq 2^r$. $\square$

The next lemma gives an upper bound on $h(d)$ that holds once d is small enough.

Lemma 17. *There is an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, and all $d \leq d_0$,*

$$h(d) \leq \left(1 + \frac{1}{\ln n}\right) \cdot (p_1 + p_{n-1}) \cdot \frac{d}{n}. \quad (60)$$

Proof. We start by showing that there is a constant $C > 0$ such that

$$h(d) \leq (p_1 + p_{n-1}) \frac{d}{n} + C(d/n)^2 \quad (61)$$

Note that $B(n, d, 1) = d/n$. We have from Lemma 16

$$h(d) = (p_1 + p_{n-1}) \frac{d}{n} + \sum_{r=2}^{12} (p_r + p_{n-r}) B(n, d, r) + \left(\frac{d}{n}\right)^2. \quad (62)$$

If there is some constant C' such that $B(n, d, r) \leq C'(d/n)^2$ for all $2 \leq r \leq 12$, then $C := C' + 1$ will do. For $r = 2$, we calculate $B(n, d, 2) = 2 \binom{d}{2} / \binom{n}{2} \leq 2(d/n)^2$, as $n > d$. Let $3 \leq \tilde{r} \leq 12$. We have

$$B(n, d, \tilde{r}) \leq \sum_{i=\lceil \tilde{r}/2 \rceil}^{\tilde{r}} (2i - \tilde{r}) \frac{\binom{d}{i} \binom{n-d}{\tilde{r}-i}}{\binom{n}{\tilde{r}}} < \tilde{r}^2 \cdot \max_{\lceil \tilde{r}/2 \rceil \leq i \leq \tilde{r}} \frac{\binom{d}{i} \binom{n-d}{\tilde{r}-i}}{\binom{n}{\tilde{r}}} \quad (63)$$

Using the inequalities $\left(\frac{n}{k}\right)^k \leq \binom{n}{k} \leq \left(\frac{ne}{k}\right)^k$, we get

$$\frac{\binom{d}{i} \binom{n-d}{\tilde{r}-i}}{\binom{n}{\tilde{r}}} \leq \frac{\frac{d^i e^i}{i^i} \cdot \frac{(n-d)^{\tilde{r}-i} e^{\tilde{r}-i}}{(\tilde{r}-i)^{\tilde{r}-i}}}{\frac{n^{\tilde{r}}}{\tilde{r}^{\tilde{r}}}} \leq e^{\tilde{r}} \cdot \tilde{r}^{\tilde{r}} \cdot \left(\frac{d}{n}\right)^i. \quad (64)$$

As $d/n \leq 1$ and $\lceil \tilde{r}/2 \rceil \geq \lceil 3/2 \rceil = 2$, we have that the maximum in (63) is at most $e^{\tilde{r}} \cdot \tilde{r}^{\tilde{r}}$, so $C' = e^{12} 12^{14}$ works.

We define $n_0 := e^C$. By our assumption that $d \leq d_0$, we have

$$\begin{aligned} h(d) &\leq (p_1 + p_{n-1}) \cdot \frac{d}{n} + \frac{C}{\ln^2 n} \cdot (p_1 + p_{n-1}) \cdot \frac{d}{n} \\ &\leq \left(1 + \frac{1}{\ln n}\right) \cdot (p_1 + p_{n-1}) \cdot \frac{d}{n}, \end{aligned} \quad (65)$$

where the last inequality holds for all $n \geq n_0$. $\square$

As we show in the following lemma, the function h is indeed an upper bound for the change in the potential. We need to show that, under our assumptions, the expected change in the potential conditioning on $d(v_t) = d + \Delta$ is maximal if $\Delta = 0$. The proof relies on a case distinction to deal with different ranges of d and Δ . Depending on the case, we use an additive Chernoff bound, a multiplicative Chernoff bound, or Lemma 17.

Lemma 18. *There is an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, any static unary unbiased algorithm with flip distribution $\mathcal{D}$ such that $p_1 + p_{n-1} = n^{-o(1)}$, and all $d \leq d_0$, it holds*

$$\mathbb{E}[X_t - X_{t+1} \mid X_t = d] \leq h(d). \quad (66)$$

Proof. We start by defining some notation

$$h_d(\tilde{d}) := \mathbb{E}[X_t - X_{t+1} \mid \{X_t = d\} \wedge \{d(v_t) = \tilde{d}\}]. \quad (67)$$

Note that

$$\mathbb{E}[X_t - X_{t+1} \mid X_t = d] \leq \max_{d \leq \tilde{d} \leq n/2} h_d(\tilde{d}). \quad (68)$$

If $d(v_t) = d$, it follows from Lemma 15 that the drift is exactly $h(d)$, i.e., we have $h_d(d) = h(d)$. So it remains to show that the above maximum is attained at $\tilde{d} = d$. We write $\Delta := \tilde{d} - d$. We need to show for all $n/2 - d \geq \Delta \geq 1$ that

$$h_d(d + \Delta) \leq h(d). \quad (69)$$

We have

$$h(d) \geq (p_1 + p_{n-1}) \frac{d}{n} \geq \frac{d}{n^{1+o(1)}}. \quad (70)$$

We observe that to make progress beyond the best-so-far search point, the algorithm needs to flip at least $\Delta + 1$, and at most $n - \Delta - 1$ bits. Together with a similar calculation as in the proof of Lemma 15, this yields

$$h_d(d + \Delta) = \sum_{r=\Delta+1}^{n-\Delta-1} (p_r + p_{n-r}) B_d(n, d + \Delta, r), \quad (71)$$

where

$$B_d(n, d + \Delta, r) = \sum_{i=\max\{\lceil \frac{r+\Delta}{2} \rceil, r+d+\Delta-n\}}^{\min\{d+\Delta, r\}} (2i - r - \Delta) \frac{\binom{d+\Delta}{i} \binom{n-d-\Delta}{r-i}}{\binom{n}{r}}. \quad (72)$$

We remark that $B_d(n, d + \Delta, r) \leq B(n, d + \Delta, r)$ and $h_d(d + \Delta) \leq h(d + \Delta)$.

As in the proof of Lemma 16, the fraction in equation (72) is equal to the probability of a hypergeometric random variable X with parameters $n, r, d + \Delta$ being equal to i . Let X be such a random variable. It holds $\mathbb{E}[X] = \frac{r(d+\Delta)}{n} \leq \frac{r}{2}$.

First, we consider the case where $\Delta \geq n/12$. Using the same argument as for equation (59), and then applying Theorem 5, we get

$$\begin{aligned} B_d(n, d + \Delta, r) &\leq r \Pr \left[X \geq \frac{r + \Delta}{2} \right] \leq r \Pr \left[X \geq \mathbb{E}[X] + \frac{\Delta}{2} \right] \\ &\leq r \exp \left(-\frac{\Delta^2}{2n} \right) \leq n \exp \left(-\frac{n}{288} \right). \end{aligned} \quad (73)$$

We have

$$h_d(d + \Delta) \leq n^2 e^{-\Omega(n)} = o(h(d)). \quad (74)$$

Next, we consider the case $n/12 > \Delta \geq \ln^2 n$. We proceed as in the proof of Lemma 16, applying Theorem 6.

$$\begin{aligned} B(n, d + \Delta, r) &\leq r \Pr \left[X \geq \frac{r}{2} \right] = r \Pr \left[X \geq \frac{n}{2(d + \Delta)} \mathbb{E}[X] \right] \\ &\leq r \left(\frac{2(d + \Delta)e}{n} \right)^{r/2} \leq r \left(\frac{e}{3} \right)^{r/2}, \end{aligned} \quad (75)$$

where we used our assumptions $d, \Delta \leq n/12$ for the last inequality. In particular, we have for $r > \ln^2 n$

$$B(n, d + \Delta, r) = O(n^{-\Omega(\ln n)}). \quad (76)$$

We have

$$h_d(d + \Delta) = n O(n^{-\Omega(\ln n)}) = o(h(d)). \quad (77)$$

It remains the case where $1 \leq \Delta < \ln^2 n$. We distinguish two subcases: (a) $d < n^{1/4} - \ln^2 n$, and (b) $d \geq n^{1/4} - \ln^2 n$.

In case (a), we have by the same argument as for (75)

$$B(n, d + \Delta, r) \leq r \left(\frac{2(d + \Delta)e}{n} \right)^{r/2} < n \left(\frac{2e}{n^{-3/4}} \right)^{r/2}. \quad (78)$$

For $r \geq 9$, we thus get $B(n, d + \Delta, r) \leq \Theta(n^{-2.375})$. For $3 \leq r \leq 8$, we get $B(n, d + \Delta, r) \leq \Theta(n^{-9/8})$. Furthermore, we have for $r = 2$, where only $\Delta = 1$ is relevant

$$B_d(n, d + 1, 2) = \frac{\binom{d+1}{2}}{\binom{n}{2}} \leq \frac{n^{2/4}e^2}{n^2} \leq \Theta(n^{-3/2}). \quad (79)$$

We compute

$$h_d(d + \Delta) \leq \Theta(n^{-3/2}) + 6 \cdot \Theta(n^{-9/8}) + n\Theta(n^{-2.375}) = \Theta(n^{-9/8}) = o(h(d)). \quad (80)$$

In case (b), it holds by Lemma 17, and picking n_0 large enough such that $n_0^{1/4} \ln n_0 - n_0^{1/4} - \ln^3 n_0 > 0$, which implies $d(\ln n - 1) > \ln^2 n$,

$$\begin{aligned} h_d(d + \Delta) &\leq h(d + \Delta) - (p_1 + p_{n-1}) \frac{d + \Delta}{n} \leq \frac{1}{\ln n} \cdot (p_1 + p_{n-1}) \frac{d + \Delta}{n} \\ &\leq \frac{1}{\ln n} \cdot (p_1 + p_{n-1}) \frac{d + \ln^2 n}{n} < \frac{1}{\ln n} \cdot (p_1 + p_{n-1}) \frac{d + d(\ln n - 1)}{n} \\ &\leq h(d), \end{aligned} \quad (81)$$

concluding the proof of this lemma. $\square$

Finally, we show that there is an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, the function h is monotonically increasing.

Lemma 19. *There is an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, the function h is monotonically increasing.*

Proof. Let $1 \leq d \leq d_0 - 1$. Take $n_0 = e^2$. Then we have $d/n \leq d_0/n \leq 1/\ln^2 n \leq 1/4$. We will show that $h(d) \leq h(d + 1)$. Recall that

$$h(d) = \sum_{r=1}^{n-1} (p_r + p_{n-r}) B(n, d, r). \quad (82)$$

Let $1 \leq r \leq n - 1$. We will show that

$$B(n, d, r) \leq B(n, d + 1, r). \quad (83)$$

Recall the definition

$$B(n, d, r) = \sum_{i=\max\{\lceil r/2 \rceil, r+d-n\}}^{\min\{d, r\}} (2i - r) \frac{\binom{d}{i} \binom{n-d}{r-i}}{\binom{n}{r}}. \quad (84)$$

The range of this sum for $B(n, d, r)$ is contained in the range of this sum for $B(n, d + 1, r)$, as otherwise, we would have $r + d + 1 - n > \lceil r/2 \rceil \geq r/2$. However, $r + d + 1 - n \leq r + n/4 - n < r - 3r/4 = r/4$, which yields a contradiction.

It remains to show that for all $\max\{\lceil r/2 \rceil, r + d - n\} \leq i \leq \min\{d, r\}$, it holds

$$(2i - r) \frac{\binom{d}{i} \binom{n-d}{r-i}}{\binom{n}{r}} \leq (2i - r) \frac{\binom{d+1}{i} \binom{n-d-1}{r-i}}{\binom{n}{r}}, \quad (85)$$

or equivalently,

$$\frac{d!}{i!(d-i)!} \cdot \frac{(n-d)!}{(r-i)!(n-d-r+i)!} \leq \frac{(d+1)!}{i!(d+1-i)!} \cdot \frac{(n-d-1)!}{(r-i)!(n-d-1-r+i)!} \quad (86)$$

Equation (86) can be rewritten as

$$\frac{n-d}{n-d-r+i} \leq \frac{d+1}{d+1-i}. \quad (87)$$

Multiplying by the denominators (which are always positive) and expanding the product, this is equivalent to

$$dn - d^2 + n - d - in + di \leq dn + n - d^2 - d - dr - r + di + i. \quad (88)$$

Finally, we see that this is the same as

$$dr + r \leq in + i, \quad (89)$$

or equivalently

$$r(d+1) \leq i(n+1). \quad (90)$$

This inequality is indeed true, as

$$i(n+1) \geq \frac{r}{2}(n+1) = r\left(\frac{n}{2} + \frac{1}{2}\right) \geq r\left(\frac{n}{4} + 1\right) \geq r(d+1), \quad (91)$$

concluding the proof of this lemma, as $h(d_0) \leq n = h(d_0 + 1)$, and h is equal to n on $[d_0 + 1, n]$. $\square$

With this lemma at hand and Lemma 13 from [DDY20], which bounds the probability to make large jumps, we can then show the following theorem, using very similar computations as those that were used in [DDY20].

Lemma 20 (Lemma 13, [DDY20]). *There exists an $n_0 \in \mathbb{N}$ such that for all $n \geq n_0$, for all $r \in [0, n]$, and for all $x \in \{0, 1\}^n$, it holds that*

$$\Pr[d(\text{flip}_r(x)) \geq \tilde{c}(d(x))] \geq 1 - n^{-4/3} \ln^7 n, \quad (92)$$

where

$$\tilde{c} : [n] \rightarrow [0, n], i \mapsto \tilde{c}(i) := \begin{cases} i - \sqrt{n} \ln n, & \text{if } i \geq n/6 \\ i - \ln^2 n, & \text{if } n^{1/3} \leq i < n/6 \\ i - 1, & \text{if } i < n^{1/3}. \end{cases} \quad (93)$$

Theorem 21. *The expected runtime of any static unary unbiased algorithm with flip distribution $\mathcal{D}$ satisfying $p_1 + p_{n-1} = n^{-o(1)}$ on any function $f : \{0, 1\}^n \rightarrow \mathbb{R}$ with unique global optimum is at least*

$$\sum_{d=1}^{d_0} \frac{1}{h(d)} - o(n). \quad (94)$$

Proof of Theorem 21. Assume that n is large enough so that all the lemmas above apply. Note that the time T_A that a unary unbiased algorithm takes to optimize f is at least as large as the hitting time T of 0 of the potential X_t . Every algorithm A has to generate its initial search point $x^{(0)}$ uniformly at random. We have $\mathbb{E}[d(x^{(0)})] = \mathbb{E}[X_0] = n/2$. By the bounds of Theorem 5, $\Pr[X_0 < n/4] \leq 2e^{-n/8}$. For now, we will assume that $X_0 \geq n/4$.

Let $c : [n] \rightarrow [0, n]$, $i \mapsto c(i) := \min\{\tilde{c}(j) \mid j \geq i\}$. Furthermore, c is monotonically increasing. It holds $c(i) \leq \tilde{c}(i)$ for all $i \in [n]$ by definition. By Lemma 20, for all $r \in [0, n]$, $x \in \{0, 1\}^n$, and $d' \leq d(x)$,

$$\Pr[d(\text{flip}_r(x)) \geq c(d')] \geq \Pr[d(\text{flip}_r(x)) \geq \tilde{c}(d')] \geq 1 - n^{-4/3} \ln^7 n. \quad (95)$$

By the law of total probability, it follows

$$\Pr[X_{t+1} \geq c(X_t)] \geq \sum_{r=0}^n p_r \Pr[d(\text{flip}_r(v_t)) \geq c(d(v_t))] \geq 1 - \underbrace{n^{-4/3} \ln^7 n}_{=:p}. \quad (96)$$

Together with Lemma 18, we see that all conditions of Theorem 4 are satisfied for $(X_t)_{t \geq 0}$, h , c , and p . It holds $\mu(x) = \max\{i \mid c(i) \leq x\} = \max\{i \mid \min\{\tilde{c}(j) \mid j \geq i\} \leq x\} = \max\{i \mid \tilde{c}(i) \leq x\}$, yielding

$$\mu(i) = \begin{cases} i+1 & \text{for } i < n^{1/3} - \ln^2 n, \\ i + \ln^2 n & \text{for } n^{1/3} - \ln^2 n \leq i < n/6 - \sqrt{n} \ln n, \\ i + \sqrt{n} \ln n & \text{for } n/6 - \sqrt{n} \ln n \leq i < n/2 - \sqrt{n} \ln n, \\ \lfloor n/2 \rfloor & \text{for } n/2 - \sqrt{n} \ln n \leq i < n/2. \end{cases} \quad (97)$$

Recall that $g(x) = \sum_{i=0}^{x-1} 1/h(\mu(x))$. By Theorem 4, we have

$$\mathbb{E}[T \mid X_0] \geq g(X_0) - \frac{g^2(X_0)p}{1 + g(X_0)p}. \quad (98)$$

First, we will bound $\frac{g^2(X_0)p}{1 + g(X_0)p}$. We have $h(d) \geq (p_1 + p_{n-1}) \cdot B(n, d, 1) = (p_1 + p_{n-1}) \cdot \frac{d}{n}$. As h is monotonically increasing, and $\mu(i) \geq i$, we have $h(x) \leq h(\mu(x))$. Hence,

$$g(X_0) \leq \sum_{i=1}^{x-1} \frac{1}{h(x)} \leq \frac{1}{p_1 + p_{n-1}} \sum_{i=1}^{x-1} \frac{n}{x} = \frac{1}{p_1 + p_{n-1}} O(n \ln n) = n^{1+o(1)}. \quad (99)$$

It follows that

$$\frac{g^2(X_0)p}{1 + g(X_0)p} = o(n). \quad (100)$$

Next, we bound $g(X_0)$ below. As h is monotonic, and all summands are positive,

$$\begin{aligned} g(X_0) &= \sum_{x=0}^{X_0-1} \frac{1}{h(\mu(x))} \geq \sum_{i=1}^{n^{1/3} - \ln^2 n} \frac{1}{h(x)} + \sum_{x=n^{1/3}}^{n/6 - \sqrt{n} \ln n + \ln^2 n} \frac{1}{h(x)} + \sum_{x=n/6}^{X_0} \frac{1}{h(x)} \\ &\geq \sum_{x=1}^{X_0} \frac{1}{h(x)} - \frac{\ln^2 n}{h(n^{1/3} - \ln^2 n)} - \frac{\sqrt{n} \ln n - \ln^2 n}{h(n/6 - \sqrt{n} \ln n + \ln^2 n)}. \end{aligned} \quad (101)$$

We apply $h(x) \geq (p_1 + p_{n-1}) \frac{x}{n}$ again and get

$$\begin{aligned} g(X_0) &\geq \sum_{x=1}^{X_0} \frac{1}{h(x)} - \frac{1}{p_1 + p_{n-1}} \cdot \left(\frac{n \ln^2 n}{n^{1/3} - \ln^2 n} + \frac{n(\sqrt{n} \ln n - \ln^2 n)}{n/6 - \sqrt{n} \ln n + \ln^2 n} \right) \\ &\geq \sum_{x=1}^{X_0} \frac{1}{h(x)} - \frac{1}{p_1 + p_{n-1}} \cdot \Theta(n^{2/3} \ln^2 n) \geq \sum_{x=1}^{X_0} \frac{1}{h(x)} - o(n). \end{aligned} \quad (102)$$

Finally, we can conclude using the law of total expectation

$$\mathbb{E}[T_A \mid X_0] \geq \left(1 - 2e^{-n/8}\right) \cdot \left(\sum_{x=1}^{n/4} \frac{1}{h(x)} - o(n) - o(n)\right) \geq \sum_{x=1}^{d_0} \frac{1}{h(x)} - o(n), \quad (103)$$

where we used that $n/4 \geq d_0$ for n large enough. $\square$

With this statement at hand, we can finally prove Theorem 14.

Proof of Theorem 14. By Lemma 17, we have for all $1 \leq d \leq d_0 = n \frac{p_1 + p_{n-1}}{\ln^2 n}$,

$$\frac{1}{h(d)} \geq \frac{n}{(p_1 + p_{n-1})d} - \frac{\frac{1}{\ln n}}{1 + \frac{1}{\ln n}} \frac{n}{(p_1 + p_{n-1})d} = (1 - o(1)) \frac{n}{(p_1 + p_{n-1})d}. \quad (104)$$

Applying Theorem 21 yields

$$\mathbb{E}[T] \geq \left((1 - o(1)) \frac{n}{(p_1 + p_{n-1})} \sum_{d=1}^{d_0} \frac{1}{d} \right) - o(n) = (1 \pm o(1)) \frac{1}{(p_1 + p_{n-1})} n \ln n. \quad (105)$$

$\square$

4.1 On p_{n-1}

The following lemma shows that the term p_{n-1} in Theorem 14 is really necessary. In particular, there are (artificial) functions on which a unary unbiased $(1+1)$ algorithm with $p_1 = 0$ can be as efficient as random local search (RLS) on ONEMAX.

Lemma 22. *Let n be even. Let RLS be the $(1+1)$ -EA $_{\mathcal{D}}$ with $p_1 = 1$ and $p_i = 0$ for $i \neq 1$, and let $\mathcal{A}$ be the $(1+1)$ -EA $_{\mathcal{D}}$ with $p_{n-1} = 1$ and $p_i = 0$ for $i \neq n-1$. Let $f : \{0, 1\}^n \rightarrow \mathbb{R}$ be defined via*

$$f(x) := \begin{cases} \text{OM}(x), & \text{if } \text{OM}(x) \text{ is even,} \\ n - \text{OM}(x), & \text{if } \text{OM}(x) \text{ is odd.} \end{cases}$$

Let $T^{\text{RLS}}(\text{OM})$ be the runtime of RLS on ONEMAX, and let $T^{\mathcal{A}}(f)$ be the runtime of $\mathcal{A}$ on f . Then $T^{\text{RLS}}(\text{OM})$ and $T^{\mathcal{A}}(f)$ follow the same distribution, i.e., for all $T \in \mathbb{N}$,

$$\Pr[T^{\text{RLS}}(\text{OM}) = T] = \Pr[T^{\mathcal{A}}(f) = T]. \quad (106)$$

In particular, $T^{\mathcal{A}}(f) = (1 \pm o(1))n \ln n$ in expectation and with high probability.

Proof. Let X_t^{RLS} and $X_t^{\mathcal{A}}$ be the fitness of RLS and $\mathcal{A}$ after t iterations, respectively. We will show by induction over t that those two random variables follow the same distribution.

Note that f is obtained from ONEMAX by swapping the fitness levels k and $n - k$ if k is odd. In particular, the number of search points of fitness k does not change. Since the initial search point is chosen uniformly at random, therefore X_0^{RLS} and $X_0^{\mathcal{A}}$ follow the same distribution.

Now let $t \geq 0$ and assume that X_t^{RLS} and $X_t^{\mathcal{A}}$ follow the same distribution. We claim that for every $k, k' \in [0..n]$, we have

$$\Pr[X_{t+1}^{\text{RLS}} = k' \mid X_t^{\text{RLS}} = k] = \Pr[X_{t+1}^{\mathcal{A}} = k' \mid X_t^{\mathcal{A}} = k]. \quad (107)$$

Note that this implies that X_{t+1}^{RLS} and $X_{t+1}^{\mathcal{A}}$ follow the same distribution since then

$$\begin{aligned}\Pr[X_{t+1}^{\text{RLS}} = k'] &= \sum_{k \in [0, n]} \Pr[X_t^{\text{RLS}} = k] \cdot \Pr[X_{t+1}^{\text{RLS}} = k' \mid X_t^{\text{RLS}} = k] \\ &= \sum_{k \in [0, n]} \Pr[X_t^{\mathcal{A}} = k] \cdot \Pr[X_{t+1}^{\mathcal{A}} = k' \mid X_t^{\mathcal{A}} = k] \\ &= \Pr[X_{t+1}^{\mathcal{A}} = k'].\end{aligned}\tag{108}$$

Moreover, since this implies that X_t^{RLS} and $X_t^{\mathcal{A}}$ follow the same distribution for all t , by

$$\Pr[T^{\text{RLS}}(\text{OM}) > T] = \Pr[X_T^{\text{RLS}} < n] = \Pr[X_T^{\mathcal{A}} < n] = \Pr[T^{\mathcal{A}}(f) > T],\tag{109}$$

it also implies the lemma. So it remains to show (107).

For RLS it is obvious that the left hand side of (107) is zero for all $k' \in [0..n] \setminus \{k, k+1\}$. Let us assume that $\mathcal{A}$ is in a search point x such that $X_t^{\mathcal{A}} = k$. The algorithm $\mathcal{A}$ creates offspring y by randomly flipping $n-1$ positions of x . This can be equivalently expressed by first flipping all n positions, and then flipping back a uniformly random position. Flipping all n positions yields the antipodal search point x' with $\text{OM}(x') = n - \text{OM}(x)$. Since y is obtained from x' by flipping exactly one bit, it satisfies $\text{OM}(y) = n - \text{OM}(x) \pm 1$. Since n is even, this implies that $\text{OM}(x')$ and $\text{OM}(x)$ are either both odd or both even. In either case, $f(x') = n - f(x)$ and thus $f(y) = f(x) \pm 1$ by definition of f . Since $\mathcal{A}$ is elitist, it will reject any offspring of fitness $f(x) - 1$, so it accepts y if and only if $f(y) = f(x) + 1$. In particular, this means that the right hand side of (107) is zero for all $k' \in [0..n] \setminus \{k, k+1\}$, as required.

For the remaining values $k' \in \{k, k+1\}$, it suffices to show equality for one of them, since the left and right hand side of (107) both sum up to one if summed over all k' . If $\text{OM}(x)$ is even, then the offspring is fitter if and only if the bit that is *not* flipped is a zero-bit, which happens with probability $(n - \text{OM}(x))/n = (n - k)/n$. If $\text{OM}(x)$ is odd, then the offspring is fitter if and only if the bit that is not flipped is a one-bit, which happens with probability $\text{OM}(x)/n = (n - k)/n$. So in either case, $\Pr[X_{t+1}^{\mathcal{A}} = k+1 \mid X_t^{\mathcal{A}} = k] = (n - k)/n$, which is the same as the probability for RLS. This concludes the proof of (107) and of the lemma. $\square$

The following theorem strengthens Theorem 14 for the $(1+1)\text{-EA}_{\mathcal{D}}$ on linear functions. It says that in this case, p_{n-1} does not help to improve the asymptotic expected runtime.

Theorem 23. *Consider the $(1+1)\text{-EA}_{\mathcal{D}}$ with distribution $\mathcal{D} = (p_0, p_1, \dots, p_n)$ such that $p_1 = n^{-o(1)}$. The expected runtime on any linear function on $\{0, 1\}^n$ is at least*

$$(1 - o(1)) \frac{1}{p_1} n \ln n.\tag{110}$$

Proof. Recall that the weights w_i of f are positive and sorted. We may assume that the smallest weight is $w_1 = 1$, since we can multiply all weights with the same constant factor without changing the fitness landscape. Moreover, for linear functions it is slightly more convenient to work with minimization instead of maximization. Both versions are equivalent, so we may assume that f is minimized. Finally, if $w_n > \sum_{i=1}^{n-1} w_i + 1$ then replacing w_n by $\sum_{i=1}^{n-1} w_i + 1$ does not change the fitness landscape since in either case all search points x with $x_n = 1$ have higher objective than all search points with $x_n = 0$. Hence we may assume $w_n \leq \sum_{i=1}^{n-1} w_i + 1$. Writing $W := \sum_{i=1}^n w_i$ for the total weight, this implies $2w_n \leq W + 1 < \frac{3}{2}W$, and thus $w_n < \frac{3}{4}W$.

Fix some t , and let $q_{i,t} := \Pr[x_i^{(t)} = 1]$ be the probability that the i -th bit is a one-bit in generation t . Then a classical result by Jägersküpfer [Jäg08] says that $q_{1,t} \geq \dots \geq q_{n,t}$. Jägersküpfer proved it for the $(1+1)\text{-EA}$ with standard bit mutation, but the only ingredient in the proof was that for all $i, j \in [n]$, if we condition on the set of flips in $[n] \setminus \{i, j\}$ then positions i and j have the same probability of being flipped. This is true for all unbiased mutation operators, so Jägersküpfer's result holds for the $(1+1)\text{-EA}_{\mathcal{D}}$ as well. Moreover, the proof shows inductively for all times that for any substring $\tilde{x}$ on the positions $[n] \setminus \{i, j\}$, the combination " $x_i = 0, x_j = 1, \tilde{x}$ " is more likely than the combination " $x_i = 1, x_j = 0, \tilde{x}$ " if $i < j$. Since both options have the same number of one-bits, and since other options (with $x_i = x_j$) contribute equally to $q_{i,t}$ and $q_{j,t}$, it was already observed in [LS15] that $q_{i,t} \geq q_{j,t}$ still holds

if we condition on the number of one-bits $\text{OM}(x^{(t)})$ at time t . Moreover, the statement also still holds if we replace t by the hitting time $T = T(d) = \min\{t \geq 0 \mid \text{OM}(x^{(t)}) \leq d\}$, so we have $q_{1,T} \geq \dots \geq q_{n,T}$.

We choose $T = T(d)$ for $d = n/\ln n$. Then we have $\sum_{i \in [n]} q_{i,T} = \mathbb{E}[\text{OM}(x^T)] \leq d$ by definition of T , and hence

$$\mathbb{E}[f(x^{(T)})] = \sum_{i \in [n]} w_i \cdot q_{i,T} \leq \frac{(\sum_{i \in [n]} w_i) \cdot (\sum_{i \in [n]} q_{i,T})}{n} \leq \frac{Wd}{n} = \frac{W}{\ln n}, \quad (111)$$

where the second step is Chebyshev's sum inequality, since w_i and $q_{i,T}$ are sorted opposingly.

By Markov's inequality, at time T we have w.h.p. $f(x^{(T)}) \leq W/8$. In the following we will condition on this event. We claim that then after time T , any offspring obtained by an $(n-1)$ -bit flip is rejected. To see this, consider any x with $f(x) \leq W/8$. Any offspring y that is obtained from x by an $(n-1)$ -bit flip has objective $f(y) \geq W - W/8 - w_n$, because the antipodal point of x has objective $W - f(x) \geq W - W/8$, and flipping back a bit can decrease the objective by at most $w_n < \frac{3}{4}W$. Hence, $f(y) \geq W - W/8 - w_n > W/8$. Therefore, the offspring y has higher (worse) objective, and is rejected. Hence, once the algorithm reaches objective at most $W/8$, all offspring obtained from $(n-1)$ -bit flips are rejected. In other words mutations of $n-1$ bits are idle steps. This means that after time T , the $(1+1)\text{-EA}_{\mathcal{D}}$ behaves as the $(1+1)\text{-EA}_{\mathcal{D}'}$, where we define $\mathcal{D}' = (p'_0, p'_1, \dots, p'_n)$ by

$$p'_i := \begin{cases} 0 & \text{if } i = n-1, \\ p_0 + p_{n-1} & \text{if } i = 0, \\ p_i & \text{otherwise.} \end{cases} \quad (112)$$

At time T w.h.p. we have $\text{OM}(x^{(T)}) \geq d - \ln^2 n$, which follows from [DDY20, Lemma 13]. By Theorem 14, the $(1+1)\text{-EA}_{\mathcal{D}'}$ needs in expectation at least $(1 - o(1)) \frac{1}{p_1 + p_{n-1}} n \ln n = (1 - o(1)) \frac{1}{p_1} n \ln n$ steps to find the optimum from level $d - \ln^2 n$, and hence the $(1+1)\text{-EA}_{\mathcal{D}}$ needs the same time. Note that we proved the lower bound conditional on w.h.p. events, but this just adds another $(1 - o(1))$ factor for the unconditional expectation. $\square$

4.2 No Stochastic Domination

Earlier work [Sud13, DJW10, Wit13] used stochastic domination arguments (cf. [Doe19]) to prove lower bounds. In particular, Witt proved his lower bound by showing that ONEMAX is the easiest function for the $(1+1)\text{-EA}$ with standard bit mutation of arbitrary mutation rate $p \leq 1/2$ [Wit13]. The key ingredient was Lemma 6.1 in [Wit13], which considered offspring y and y' that are created from x and x' respectively by standard bit mutation with mutation rate $p \leq 1/2$. For minimization, if $\text{OM}(x) \leq \text{OM}(x')$ then the lemma states $\Pr[\text{OM}(y) \leq k] \geq \Pr[\text{OM}(y') \leq k]$ for all $k \in [0, n]$. So it is easier to reach OM-level at most k when starting with a parent of smaller OM-value. This lemma implies on the one hand that ONEMAX is the easiest function for standard bit mutation, but also that elitist selection is optimal in this situation: the $(1+1)\text{-EA}$ with mutation rate $p \leq 1/2$ is the fastest algorithm on ONEMAX among all unary algorithms using standard bit mutation with mutation rate $p \leq 1/2$.

However, Witt's lemma does not hold for general unbiased mutation operators. In particular, being closer to the optimum does not mean that we have a higher chance of finding the optimum in the next step. Consider the case where the algorithm flips one bit with probability $p_1 = n^{-2}$ and two bits with probability $p_2 = 1 - p_1 = 1 - n^{-2}$. The probability of finding the optimum from a search point in Hamming distance one from the optimum is $p_1/n = n^{-3}$, whereas the probability of finding the optimum from a search point in Hamming distance two from the optimum is $p_2/\binom{n}{2} = \Theta(n^{-2})$, which is much larger.

Even worse, let us consider the time T_d to find the optimum on ONEMAX if we start in Hamming distance d . For $d = 1$ we have $\mathbb{E}[T_1] = n/p_1 = \Theta(n^3)$. For $d = 2$, the probability of making an improvement is $p_{\text{imp}} = p_1 \cdot 2/n + p_2/\binom{n}{2} = \Theta(n^{-2})$. Hence, conditional on making an improvement, the algorithm improves by one with probability $(p_1 \cdot 2/n)/p_{\text{imp}} = \Theta(n^{-1})$. Therefore, we need to wait in expectation $1/p_{\text{imp}} = \Theta(n^2)$ rounds for an improvement, and with probability $\Theta(n^{-1})$ we improve only by one and need to wait another T_1 rounds for reaching the optimum. Hence,

$$\mathbb{E}[T_2] = \Theta(n^2) + \Theta(n^{-1}) \cdot \mathbb{E}[T_1] = \Theta(n^2), \quad (113)$$

which is asymptotically smaller than $\mathbb{E}[T_1] = \Theta(n^3)$. So the expected time $\mathbb{E}[T_d]$ is not monotone in d , and can be asymptotically smaller if we start further away from the optimum.

Turning this example around, we can construct a situation where ONEMAX is not the easiest function. Consider an algorithm with $p_1 = n^{-3}$, $p_2 = n^{-1}$ and $p_3 = 1 - p_1 - p_2 = 1 - o(1)$, starting in the string $x = (01 \dots 1)$ where all but the first bit are optimized. On ONEMAX, it needs to wait for a one-bit flip, which takes time $n/p_1 = \Theta(n^4)$. But if the fitness function is $f(x) := 3x_1 + \sum_{i=2}^n x_i$, then the algorithm accepts any mutation flipping two or three bits if it involves x_1 . Conditional on flipping x_1 , a two-bit flip has only probability $O(n^{-1})$ since $p_3/p_2 = \Theta(n)$. In that case (an improving two-bit flip) the algorithm jumps to another neighbour of the optimum and needs to wait $n/p_1 = O(n^4)$ rounds for the right one-bit flip. This contributes $O(n^{-1} \cdot n^4) = O(n^3)$ to the expectation. However, in the more likely case of a three-bit flip, the algorithm jumps to a search point in distance two from the optimum. By a similar calculation as before, it now needs time $O(n^3)$ to find the optimum, so the expected runtime on f is $O(n^3)$, which is asymptotically faster than on ONEMAX.

Finally, the same example can be used to show that a non-elitist $(1 + 1)$ algorithm may be faster than the $(1 + 1)$ -EA $_{\mathcal{D}}$ if both use the same unbiased mutation operator. Hence, Witt’s lemma and all its consequences fail for general unbiased mutation operators. This is similar to the situation for the compact genetic algorithm cGA, for which this form of domination also does not hold [Doe21].

5 Conclusions

We have extended Witt’s result bounding the runtime of the $(1 + 1)$ -EA on linear functions to arbitrary elitist $(1 + 1)$ unary unbiased EAs and we have discussed various ways in which the requirements made in Corollary 10 and Theorem 14 are tight. In particular, we have seen that for $p_1 = n^{-\Omega(1)}$, the expected runtime can be smaller than $\frac{1}{p_1} n \ln n$ by a constant factor. When interpreted in the light of black-box complexity, our results can be seen as extensions of [DDY20] to linear functions. However, we have focused in this work on *static* mutation operators. An extension of our result to *dynamic* parameter settings would hence be a natural continuation of our work.

Another direction in which we aim to extend our results are *combinatorial* optimization problems where we suspect to see a tangible advantage of unusual unary mutation operators. For example, the optimal mutation operator for the minimum spanning tree problem (MST) is likely to satisfy $p_1 > 0$ and $p_2 > 0$. Similarly, there are functions like LEADINGONES where the optimal number of flipped bits depends on the phase of the algorithm, and none of the phases is asymptotically negligible for the runtime. In such cases, it may be interesting to see what the optimal distribution is.

Similarly, we also expect advantages of the $(1 + 1)$ -EA $_{\mathcal{D}}$ over standard $(1 + 1)$ -EAs when optimizing for average performance for problem collections with instances having different landscapes.

Acknowledgments. Our work is financially supported by ANR-22-ERCS-0003-01 project VARIATION. Duri Andrea Janett was supported by the Swiss-European Mobility Programme.

References

- [ABD21] Denis Antipov, Maxim Buzdalov, and Benjamin Doerr. Lazy parameter tuning and control: choosing all parameters randomly from a power-law distribution. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, pages 1115–1123. ACM, 2021.
- [ABD22] Denis Antipov, Maxim Buzdalov, and Benjamin Doerr. Fast mutation in crossover-based algorithms. *Algorithmica*, 84:1724–1761, 2022.
- [AD20] Denis Antipov and Benjamin Doerr. Runtime analysis of a heavy-tailed $(1 + (\lambda, \lambda))$ genetic algorithm on jump functions. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, LNCS, pages 545–559, 2020.
- [AD21] Denis Antipov and Benjamin Doerr. Precise runtime analysis for plateau functions. *ACM Transactions on Evolutionary Learning and Optimization*, 1(4):1–28, 2021.
- [AN21] Denis Antipov and Semen Naumov. The effect of non-symmetric fitness: The analysis of crossover-based algorithms on realjump functions. In *Proceedings of the 16th ACM/SIGEVO Conference on Foundations of Genetic Algorithms*, pages 1–15, 2021.
- [BBD22] Henry Bambury, Antoine Bultel, and Benjamin Doerr. An extended jump functions benchmark for the analysis of randomized search heuristics. *Algorithmica*, pages 1–32, 2022.

- [BD20] Maxim Buzdalov and Carola Doerr. Optimal mutation rates for the $(1 + \lambda)$ EA on OneMax. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, volume 12270 of *LNCS*, pages 574–587. Springer, 2020.
- [BD21a] Nathan Buskulis and Carola Doerr. Maximizing drift is not optimal for solving OneMax. *Evolutionary Computation*, 29(4):521–541, 2021.
- [BD21b] Maxim Buzdalov and Carola Doerr. Optimal static mutation strength distributions for the $(1 + \lambda)$ evolutionary algorithm on OneMax. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, pages 660–668. ACM, 2021.
- [BDM⁺21] Pauline Bennet, Carola Doerr, Antoine Moreau, Jeremy Rapin, Fabien Teytaud, and Olivier Teytaud. Nevergrad: black-box optimization platform. *ACM SIGEVOlution*, 14(1):8–15, 2021.
- [Buz22] Maxim Buzdalov. The $(1 + (\lambda, \lambda))$ genetic algorithm on the vertex cover problem: Crossover helps leaving plateaus. In *Proc. of IEEE Congress on Evolutionary Computation (CEC)*, pages 1–10. IEEE, 2022.
- [COY21a] Dogan Corus, Pietro S Oliveto, and Donya Yazdani. Automatic adaptation of hypermutation rates for multimodal optimisation. In *Proc. of ACM/SIGEVO Conference on Foundations of Genetic Algorithms (FOGA)*, pages 1–12. ACM, 2021.
- [COY21b] Dogan Corus, Pietro S Oliveto, and Donya Yazdani. Fast immune system-inspired hypermutation operators for combinatorial optimization. *IEEE Transactions on Evolutionary Computation*, 25(5):956–970, 2021.
- [DD14] Benjamin Doerr and Carola Doerr. Reducing the arity in unbiased black-box complexity. *Theoretical Computer Science*, 545:108–121, 2014.
- [DDE15] Benjamin Doerr, Carola Doerr, and Franziska Ebel. From black-box complexity to designing new genetic algorithms. *Theoretical Computer Science*, 567:87 – 104, 2015.
- [DDK15] Benjamin Doerr, Carola Doerr, and Timo Kötzing. Unbiased black-box complexities of jump functions. *Evolutionary Computation*, 23(4):641–670, 2015.
- [DDY20] Benjamin Doerr, Carola Doerr, and Jing Yang. Optimal parameter choices via precise black-box analysis. *Theoretical Computer Science*, 801:1–34, 2020.
- [DGB22] Benjamin Doerr, Yassine Ghannane, and Marouane Ibn Brahim. Runtime analysis for permutation-based evolutionary algorithms. *arXiv preprint arXiv:2207.04045*, 2022.
- [DHP22] Benjamin Doerr, Omar El Hadri, and Adrien Pinard. The $(1 + (\lambda, \lambda))$ global semo algorithm. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, pages 520–528, 2022.
- [DJK⁺11] Benjamin Doerr, Daniel Johannsen, Timo Kötzing, Per Kristian Lehre, Markus Wagner, and Carola Winzen. Faster black-box algorithms through higher arity operators. In *Proc. of Foundations of Genetic Algorithms (FOGA)*, pages 163–172, 2011.
- [DJW06] Stefan Droste, Thomas Jansen, and Ingo Wegener. Upper and lower bounds for randomized search heuristics in black-box optimization. *Theory of Computing Systems*, 39:525–544, 2006.
- [DJW10] Benjamin Doerr, Daniel Johannsen, and Carola Winzen. Drift analysis and linear functions revisited. In *Proc. of IEEE Congress on Evolutionary Computation (CEC)*, pages 1–8. IEEE, 2010.
- [DJW12] Benjamin Doerr, Daniel Johannsen, and Carola Winzen. Multiplicative drift analysis. *Algorithmica*, 64:673–697, 2012.
- [DKLW13] Benjamin Doerr, Timo Kötzing, Johannes Lengler, and Carola Winzen. Black-box complexities of combinatorial problems. *Theoretical Computer Science*, 471:84–106, 2013.
- [DLMN17] Benjamin Doerr, Huu Phuoc Le, Régis Makhlara, and Ta Duy Nguyen. Fast genetic algorithms. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, 2017.
- [DN21] Benjamin Doerr and Frank Neumann. A survey on recent progress in the theory of evolutionary algorithms for discrete optimization. *ACM Transactions on Evolutionary Learning and Optimization*, 1(4):16:1–16:43, 2021.
- [Doe19] Benjamin Doerr. Analyzing randomized search heuristics via stochastic domination. *Theoretical Computer Science*, 773:115–137, 2019.
- [Doe20] Carola Doerr. Complexity theory for discrete black-box optimization heuristics. *Theory of Evolutionary Computation: Recent Developments in Discrete Optimization*, pages 133–212, 2020.

- [Doe21] Benjamin Doerr. The runtime of the compact genetic algorithm on jump functions. *Algorithmica*, 83:1–49, 10 2021.
- [DQ22] Benjamin Doerr and Zhongdi Qu. A first runtime analysis of the nsga-ii on a multimodal problem. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, pages 399–412. Springer, 2022.
- [DR22] Benjamin Doerr and Amirhossein Rajabi. Stagnation detection meets fast mutation. *Theoretical Computer Science*, 2022.
- [DW12] Benjamin Doerr and Carola Winzen. Black-box complexity: Breaking the $O(n \log n)$ barrier of LeadingOnes. In *Proc. of Artificial Evolution (EA)*, pages 205–216. Springer, 2012.
- [DWY⁺18] Carola Doerr, Hao Wang, Furong Ye, Sander Van Rijn, and Thomas Bäck. IOHprofiler: A benchmarking and profiling tool for iterative optimization heuristics. *arXiv preprint arXiv:1810.05281*, 2018.
- [DZ21] Benjamin Doerr and Weijie Zheng. Theoretical analyses of multi-objective evolutionary algorithms on multi-modal objectives. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 35, pages 12293–12301, 2021.
- [FGQW18] Tobias Friedrich, Andreas Göbel, Francesco Quinzan, and Markus Wagner. Heavy-tailed mutation operators in single-objective combinatorial optimization. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, pages 134–145. Springer, 2018.
- [FQW18] Tobias Friedrich, Francesco Quinzan, and Markus Wagner. Escaping large deceptive basins of attraction with heavy-tailed mutation operators. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, pages 293–300. ACM, 2018.
- [Hoe63] Wassily Hoeffding. Probability inequalities for sums of bounded random variables. *Journal of the American statistical association*, 58(301):13–30, 1963.
- [Jäg08] Jens Jägersküpper. A blend of Markov-chain and drift analysis. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, pages 41–51. Springer, 2008.
- [KNSH21] Jaroslav Klapálek, Antonín Novák, Michal Sojka, and Zdeněk Hanzálek. Car racing line optimization with genetic algorithm using approximate homeomorphism. In *2021 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, pages 601–607. IEEE, 2021.
- [Len19] Johannes Lengler. A general dichotomy of evolutionary algorithms on monotone functions. *IEEE Transactions on Evolutionary Computation*, 24(6):995–1009, 2019.
- [LS15] Johannes Lengler and Nicholas Spooner. Fixed budget performance of the (1+1) EA on linear functions. In *Proce. of Foundations of Genetic Algorithms (FOGA)*, pages 52–61. ACM, 2015.
- [LS19] Per Kristian Lehre and Dirk Sudholt. Parallel black-box complexity with tail bounds. *IEEE Transactions on Evolutionary Computation*, 24(6):1010–1024, 2019.
- [LW12] Per Kristian Lehre and Carsten Witt. Black-box search by unbiased variation. *Algorithmica*, 64:623–642, 2012.
- [MB17] Vladimir Mironovich and Maxim Buzdalov. Evaluation of heavy-tailed mutation operator on maximum flow test generation problem. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO, Companion)*, pages 1423–1426, 2017.
- [NAN22] Aneta Neumann, Denis Antipov, and Frank Neumann. Coevolutionary pareto diversity optimization. In *Proc. of Genetic and Evolutionary Computation Conference (GECCO)*, pages 832–839. ACM, 2022.
- [NSN⁺22] Antonin Novak, Premysl Sucha, Matej Novotny, Richard Stec, and Zdenek Hanzalek. Scheduling jobs with normally distributed processing times on parallel machines. *European Journal of Operational Research*, 297(2):422–441, 2022.
- [NXN22] Aneta Neumann, Yue Xie, and Frank Neumann. Evolutionary algorithms for limiting the effect of uncertainty for the knapsack problem with stochastic profits. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, pages 294–307. Springer, 2022.
- [PCS22] Artem Pavlenko, Daniil Chivilikhin, and Alexander Semenov. Asynchronous evolutionary algorithm for finding backdoors in boolean satisfiability. In *Proc. of IEEE Congress on Evolutionary Computation (CEC)*, pages 1–8. IEEE, 2022.
- [QGWF21] Francesco Quinzan, Andreas Göbel, Markus Wagner, and Tobias Friedrich. Evolutionary algorithms and submodular functions: benefits of heavy-tailed mutations. *Natural*

- Computing*, pages 1–15, 2021.
- [SCP⁺21] Alexander Semenov, Daniil Chivilikhin, Artem Pavlenko, Ilya Otpuschennikov, Vladimir Ulyantsev, and Alexey Ignatiev. Evaluating the hardness of sat instances using evolutionary optimization algorithms. In *Proc. of International Conference on Principles and Practice of Constraint Programming (CP 2021)*. Schloss Dagstuhl-Leibniz-Zentrum für Informatik, 2021.
 - [Sud13] Dirk Sudholt. A new method for lower bounds on the running time of evolutionary algorithms. *IEEE Transactions on Evolutionary Computation*, 17(3):418–435, 2013.
 - [Wit13] Carsten Witt. Tight bounds on the optimization time of a randomized search heuristic on linear functions. *Combinatorics, Probability and Computing*, 22(2):294–318, 2013.
 - [Wit14] Carsten Witt. Fitness levels with tail bounds for the analysis of randomized search heuristics. *Information Processing Letters*, 114(1-2):38–41, 2014.
 - [WQT18] Mengxi Wu, Chao Qian, and Ke Tang. Dynamic mutation based pareto optimization for subset selection. In *Proc. of Intelligent Computing Methodologies*, pages 25–35. Springer, 2018.
 - [YDB19] Furong Ye, Carola Doerr, and Thomas Bäck. Interpolating local and global search by controlling the variance of standard bit mutation. In *Proc. of IEEE Congress on Evolutionary Computation (CEC)*, pages 2292–2299. IEEE, 2019.
 - [YWDB20] Furong Ye, Hao Wang, Carola Doerr, and Thomas Bäck. Benchmarking a genetic algorithm with configurable crossover probability. In *Proc. of Parallel Problem Solving from Nature (PPSN)*, pages 699–713. Springer, 2020.