



2020/06

DP

Anton Rodomanov and
Yurii Nesterov

Greedy quasi-Newton
methods with explicit
superlinear convergence



CORE

CENTER FOR
OPERATIONS RESEARCH
AND ECONOMETRICS

CORE

Voie du Roman Pays 34, L1.03.01

B-1348 Louvain-la-Neuve

Tel (32 10) 47 43 04

Email: immaq-library@uclouvain.be

[https://uclouvain.be/en/research-institutes/
lidam/core/discussion-papers.html](https://uclouvain.be/en/research-institutes/lidam/core/discussion-papers.html)

CORE DISCUSSION PAPER

2020/06

Greedy Quasi-Newton Methods with Explicit Superlinear Convergence ^{*}

Anton Rodomanov[†]

Yurii Nesterov[‡]

January 19, 2020

Abstract

In this paper, we study greedy variants of quasi-Newton methods. They are based on the updating formulas from a certain subclass of the Broyden family. In particular, this subclass includes the well-known DFP, BFGS and SR1 updates. However, in contrast to the classical quasi-Newton methods, which use the difference of successive iterates for updating the Hessian approximations, our methods apply basis vectors, greedily selected so as to maximize a certain measure of progress. For greedy quasi-Newton methods, we establish an explicit non-asymptotic bound on their rate of local superlinear convergence, which contains a contraction factor, depending on the square of the iteration counter. We also show that these methods produce Hessian approximations whose deviation from the exact Hessians linearly converges to zero.

Keywords: quasi-Newton methods, Broyden family, SR1, DFP, BFGS, superlinear convergence, local convergence, rate of convergence

^{*}The research results of this paper were obtained with support of ERC Advanced Grant 788368.

[†]Institute of Information and Communication Technologies, Electronics and Applied Mathematics (ICTEAM), Catholic University of Louvain (UCL). E-mail: anton.rodomanov@uclouvain.be.

[‡]Center for Operations Research and Econometrics (CORE), Catholic University of Louvain (UCL). E-mail: yurii.nesterov@uclouvain.be.

1 Introduction

Motivation. Quasi-Newton methods have a reputation of the most efficient numerical schemes for smooth unconstrained optimization. The main idea of these algorithms is to approximate the standard Newton method by replacing the exact Hessian with some approximation, which is updated between iterations according to special formulas. There exist numerous variants of quasi-Newton algorithms that differ mainly in the rules of updating Hessian approximations. The three most popular are the *Davidon–Fletcher–Powell (DFP)* method [1, 3], the *Broyden–Fletcher–Goldfarb–Shanno (BFGS)* method [4–8], and the *Symmetric Rank 1 (SR1)* method [1, 2]. For a general overview of the topic, see [12] and [22, Ch. 6]; also see [24] for the application of quasi-Newton methods for non-smooth optimization.

The most attractive feature of quasi-Newton methods is their *superlinear convergence*, which was first established in the 1970s [9–11]. Namely, for several standard quasi-Newton methods (such as DFP and BFGS), it was proved that the ratio of successive residuals tends to zero as the number of iterations goes to infinity. However, the authors did not obtain any explicit bounds on the corresponding *rate* of this superlinear convergence. For example, it is unknown whether the residuals convergence like $O(c^{k^2})$, where $c \in (0, 1)$ is some constant and k is the iteration counter, or $O(c^{k^3})$, or $O(k^{-k})$, or somehow else. Thus, despite the qualitative usefulness of the mentioned result, it still lacks quantitative estimates of the rate of convergence. Although many other works on quasi-Newton methods have appeared since then, to our knowledge, all of them still contain only asymptotic results (see e.g. [14–18, 20, 21, 23, 26, 29, 30]). Thus, up to now, there are still no *explicit* and *non-asymptotic* estimates of the rate of superlinear convergence of quasi-Newton methods.

In this work, we make a first step towards obtaining such estimates. We propose new quasi-Newton methods, which are based on the updating formulas from a certain subclass of the Broyden family [2]. In particular, this subclass contains the DFP, BFGS and SR1 updates. However, in contrast to the classical quasi-Newton methods, which use the difference of successive iterates for updating the Hessian approximations, our methods apply *basis vectors*, greedily selected to maximize a certain measure of progress. For greedy quasi-Newton methods, we establish an explicit non-asymptotic bound on their rate of local superlinear convergence, which contains a contraction factor, depending on the *square* of the iteration counter. We also show that these methods produce Hessian approximations whose deviation from the exact Hessians converges to zero at a *linear* rate. Note that this is not the case for the standard quasi-Newton methods, which cannot ensure any convergence in the Hessian approximation at all (see e.g. [11]).

Let us mention that the idea of using basis vectors in quasi-Newton methods goes back at least to so-called *methods of dual directions* [13], for which it is also possible to prove both local superlinear convergence of the iterates and convergence of the Hessian approximations. However, similarly to the standard quasi-Newton methods, all corresponding results are only asymptotic. In any case, despite to the fact that the greedy quasi-Newton methods, presented in this paper, are based on the same idea, their construction and analysis are significantly different.

Contents. In Section 2, we discuss a class of quasi-Newton updating rules for approximating a self-adjoint positive definite linear operator. We present a special greedy strategy for selecting an update direction, which ensures a linear convergence rate in ap-

proximating the target operator. In Section 3, we analyze greedy quasi-Newton methods, applied to the problem of minimizing a quadratic function. We show that these methods have a global linear convergence rate, comparable to that of the standard gradient descent, and also a superlinear convergence rate, which contains a contraction factor, depending on the square of the iteration counter. In Section 4, we show that similar results also hold for general nonlinear functions, provided that the starting point is chosen sufficiently close to the solution. The main difficulty here, compared to the quadratic case, is that the Hessian of the objective function is no longer constant, resulting in the necessity to apply a special *correction strategy* to keep the Hessian approximations under control. Finally, in Section 5, we present some preliminary computational results.

Notation. In what follows, $\mathbb{E}$ denotes an arbitrary n -dimensional real vector space. Its dual space, composed of all linear functionals on $\mathbb{E}$, is denoted by $\mathbb{E}^*$. The value of a linear function $s \in \mathbb{E}^*$, evaluated at a point $x \in \mathbb{E}$, is denoted by $\langle s, x \rangle$.

For a smooth function $f : \mathbb{E} \rightarrow \mathbb{R}$, we denote by $\nabla f(x)$ and $\nabla^2 f(x)$ its gradient and Hessian respectively, evaluated at a point $x \in \mathbb{E}$. Note that $\nabla f(x) \in \mathbb{E}^*$, and $\nabla^2 f(x)$ is a self-adjoint linear operator from $\mathbb{E}$ to $\mathbb{E}^*$.

The partial ordering of self-adjoint linear operators is defined in the standard way. We write $A \preceq A_1$ for $A, A_1 : \mathbb{E} \rightarrow \mathbb{E}^*$ if $\langle (A_1 - A)x, x \rangle \geq 0$ for all $x \in \mathbb{E}$, and $W \preceq W_1$ for $W, W_1 : \mathbb{E}^* \rightarrow \mathbb{E}$ if $\langle s, (W_1 - W)s \rangle \geq 0$ for all $s \in \mathbb{E}^*$.

Any self-adjoint positive definite linear operator $A : \mathbb{E} \rightarrow \mathbb{E}^*$ induces in the spaces $\mathbb{E}$ and $\mathbb{E}^*$ the following pair of conjugate Euclidean norms:

$$\begin{aligned} \|h\|_A &\stackrel{\text{def}}{=} \langle Ah, h \rangle^{1/2}, & h \in \mathbb{E}, \\ \|s\|_A^* &\stackrel{\text{def}}{=} \langle s, A^{-1}s \rangle^{1/2}, & s \in \mathbb{E}^*. \end{aligned} \tag{1.1}$$

When $A = \nabla^2 f(x)$, where $f : \mathbb{E} \rightarrow \mathbb{R}$ is a smooth function with positive definite Hessian, and $x \in \mathbb{E}$, we prefer to use notation $\|\cdot\|_x$ and $\|\cdot\|_x^*$, provided that there is no ambiguity with the reference function f .

Sometimes, in the formulas, involving products of linear operators, it is convenient to treat $x \in \mathbb{E}$ as a linear operator from $\mathbb{R}$ to $\mathbb{E}$, defined by $x\alpha = \alpha x$, and x^* as a linear operator from $\mathbb{E}^*$ to $\mathbb{R}$, defined by $x^*s = \langle s, x \rangle$. In this case, xx^* is a rank-one self-adjoint linear operator from $\mathbb{E}^*$ to $\mathbb{E}$, acting as follows:

$$(xx^*)s = \langle s, x \rangle x, \quad s \in \mathbb{E}^*.$$

Likewise, any $s \in \mathbb{E}^*$ can be treated as a linear operator from $\mathbb{R}$ to $\mathbb{E}^*$, defined by $s\alpha = \alpha s$, and s^* as a linear operator from $\mathbb{E}$ to $\mathbb{R}$, defined by $s^*x = \langle s, x \rangle$. Then, ss^* is a rank-one self-adjoint linear operator from $\mathbb{E}$ to $\mathbb{E}^*$.

For two self-adjoint linear operators $A : \mathbb{E} \rightarrow \mathbb{E}^*$ and $W : \mathbb{E}^* \rightarrow \mathbb{E}$, define

$$\langle W, A \rangle \stackrel{\text{def}}{=} \text{Trace}(WA).$$

Note that WA is a linear operator from $\mathbb{E}$ to itself, and hence its trace is well-defined (it coincides with the trace of the matrix representation of WA with respect to an arbitrary chosen basis in the space $\mathbb{E}$, and the result is independent of the particular choice of the

basis). Observe that $\langle \cdot, \cdot \rangle$ is a bilinear form, and for any $x \in \mathbb{E}$, we have

$$\langle Ax, x \rangle = \langle xx^*, A \rangle. \quad (1.2)$$

When A is invertible, we also have

$$\langle A^{-1}, A \rangle = n. \quad (1.3)$$

If the operator W is positive semidefinite, and $A \preceq A_1$ for some self-adjoint linear operator $A_1 : \mathbb{E} \rightarrow \mathbb{E}^*$, then $\langle W, A \rangle \leq \langle W, A_1 \rangle$. Similarly, if A is positive semidefinite and $W \preceq W_1$ for some self-adjoint linear operator $W_1 : \mathbb{E}^* \rightarrow \mathbb{E}$, then $\langle W, A \rangle \leq \langle W_1, A \rangle$. When A is positive definite, and $R : \mathbb{E} \rightarrow \mathbb{E}^*$ is a self-adjoint linear operator, $\langle A^{-1}, R \rangle$ equals the sum of the eigenvalues of R with respect to the operator A . In particular, if R is positive semidefinite, then all its eigenvalues with respect to A are non-negative, and the maximal one can be bounded by the trace:

$$R \preceq \langle A^{-1}, R \rangle A. \quad (1.4)$$

2 Greedy Quasi-Newton Updates

Let $A : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint positive definite linear operator. In this section, we consider a class of quasi-Newton updating rules for approximating A .

Let $G : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint linear operator, such that

$$A \preceq G, \quad (2.1)$$

and let $u \in \mathbb{E}$ be a direction. Consider the following family of updates, parameterized by a scalar $\tau \in \mathbb{R}$. If $Gu \neq Au$, define

$$\begin{aligned} \text{Broyd}_\tau(G, A, u) &\stackrel{\text{def}}{=} \tau \left[G - \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle} \right] \\ &\quad + (1 - \tau) \left[G - \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} \right]. \end{aligned} \quad (2.2)$$

Otherwise, if $Gu = Au$, define $\text{Broyd}_\tau(G, A, u) \stackrel{\text{def}}{=} G$.

Note that, for $\tau = 0$, formula (2.2) corresponds to the well-known *SR1 update*,

$$\text{SR1}(G, A, u) \stackrel{\text{def}}{=} G - \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle}, \quad (2.3)$$

and, for $\tau = 1$, it corresponds to the well-known *DFP update*:

$$\text{DFP}(G, A, u) \stackrel{\text{def}}{=} G - \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle}. \quad (2.4)$$

Thus, (2.2) describes the *Broyden family* of quasi-Newton updates (see [22, Section 6.3]), and can be written as the linear combination of DFP and SR1 updates:¹

$$\text{Broyd}_\tau(G, A, u) = \tau \text{DFP}(G, A, u) + (1 - \tau) \text{SR1}(G, A, u).$$

¹Usually, the Broyden family is defined as the linear combination of the DFP and BFGS updates. Here we use alternative (but equivalent) parametrization of this family, which is more convenient for our purposes.

Our main interest will be in the class, described by the values $\tau \in [0, 1]$, i.e. in the convex combination of the DFP and SR1 updates. Note in particular, that this subclass includes another well-known update—*BFGS*. Indeed, for

$$\tau_{\text{BFGS}} \stackrel{\text{def}}{=} \frac{\langle Au, u \rangle}{\langle Gu, u \rangle} \stackrel{(2.1)}{\in} (0, 1), \quad (2.5)$$

we have $1 - \tau_{\text{BFGS}} = \frac{\langle (G-A)u, u \rangle}{\langle Gu, u \rangle}$, and thus

$$\begin{aligned} \text{Broyd}_{\tau_{\text{BFGS}}}(G, A, u) &= G - \frac{\langle (G-A)u, u \rangle}{\langle Gu, u \rangle} \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} \\ &\quad + \frac{\langle Au, u \rangle}{\langle Gu, u \rangle} \left[-\frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle} \right] \\ &= G - \frac{(G-A)uu^*(G-A)}{\langle Gu, u \rangle} - \frac{Auu^*G + Guu^*A}{\langle Gu, u \rangle} \\ &\quad + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Gu, u \rangle} \\ &= G - \frac{Guu^*G}{\langle Gu, u \rangle} + \frac{Auu^*A}{\langle Au, u \rangle} \stackrel{\text{def}}{=} \text{BFGS}(G, A, u). \end{aligned} \quad (2.6)$$

This is the classic BFGS formula for direction u .

Let us show that the Broyden family is monotonic in the parameter τ .

Lemma 2.1 *If (2.1) holds, then, for any $u \in \mathbb{E}$ and any $\tau_1, \tau_2 \in \mathbb{R}$, such that $\tau_1 \leq \tau_2$,*

$$\text{Broyd}_{\tau_1}(G, A, u) \preceq \text{Broyd}_{\tau_2}(G, A, u).$$

Proof: Without loss of generality, suppose $Gu \neq Au$. Then

$$\begin{aligned} \text{Broyd}_{\tau}(G, A, u) &\stackrel{(2.2)}{=} G - \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} \\ &\quad + \tau \left[\frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} - \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle} \right]. \end{aligned}$$

Denote $s \stackrel{\text{def}}{=} \frac{(G-A)u}{\langle (G-A)u, u \rangle} - \frac{Au}{\langle Au, u \rangle}$. Then,

$$\begin{aligned} \langle (G-A)u, u \rangle ss^* &= \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} + \frac{\langle (G-A)u, u \rangle}{\langle Au, u \rangle} \frac{Auu^*A}{\langle Au, u \rangle} - \frac{(G-A)uu^*A + Auu^*(G-A)}{\langle Au, u \rangle} \\ &= \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} - \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle}. \end{aligned}$$

Therefore,

$$\text{Broyd}_{\tau}(G, A, u) = G - \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle} + \tau \langle (G-A)u, u \rangle ss^*.$$

The claim now follows from the fact that $\langle (G-A)u, u \rangle ss^* \succeq 0$ in view of (2.1). $\square$

Next, let us show that the relation (2.1) can be preserved after applying to G any update from the class of our interest. Moreover, each update from this class does not increase the deviation from the target operator A .

Lemma 2.2 *If, for some $\eta \geq 1$, we have*

$$A \preceq G \preceq \eta A, \quad (2.7)$$

then, for any $u \in \mathbb{E}$ and any $\tau \in [0, 1]$, we also have

$$A \preceq \text{Broyd}_\tau(G, A, u) \preceq \eta A. \quad (2.8)$$

Proof: Without loss of generality, suppose $Gu \neq Au$. Denote $G_+ \stackrel{\text{def}}{=} \text{Broyd}_\tau(G, A, u)$. Since $\tau \geq 0$, by Lemma 2.1, we have

$$G_+ \succeq \text{SR1}(G, A, u) \stackrel{(2.3)}{=} G - \frac{(G-A)uu^*(G-A)}{\langle (G-A)u, u \rangle}.$$

Let $R \stackrel{\text{def}}{=} G - A \stackrel{(2.7)}{\succeq} 0$, and let $I_{\mathbb{E}}, I_{\mathbb{E}^*}$ be the identity operators in the spaces $\mathbb{E}, \mathbb{E}^*$ respectively. Then,

$$G_+ - A \succeq R - \frac{Ru u^* R}{\langle Ru, u \rangle} = \left(I_{\mathbb{E}^*} - \frac{Ru u^*}{\langle Ru, u \rangle} \right) R \left(I_{\mathbb{E}} - \frac{uu^* R}{\langle Ru, u \rangle} \right) \succeq 0,$$

Thus, the first relation in (2.8) is proved. To prove the second relation, we apply Lemma 2.1, using that $\tau \leq 1$, to obtain

$$\begin{aligned} G_+ &\preceq \text{DFP}(G, A, u) \stackrel{(2.4)}{=} G + \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle} - \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} \\ &= \frac{Auu^*A}{\langle Au, u \rangle} + \left(I_{\mathbb{E}^*} - \frac{Auu^*}{\langle Au, u \rangle} \right) G \left(I_{\mathbb{E}} - \frac{uu^*A}{\langle Au, u \rangle} \right) \\ &\stackrel{(2.7)}{\preceq} \frac{Auu^*A}{\langle Au, u \rangle} + \eta \left(I_{\mathbb{E}^*} - \frac{Auu^*}{\langle Au, u \rangle} \right) A \left(I_{\mathbb{E}} - \frac{uu^*A}{\langle Au, u \rangle} \right) \\ &= \frac{Auu^*A}{\langle Au, u \rangle} + \eta \left(A - \frac{Auu^*A}{\langle Au, u \rangle} \right) = \eta A - (\eta - 1) \frac{Auu^*A}{\langle Au, u \rangle} \preceq \eta A. \quad \square \end{aligned}$$

Interestingly, from Lemma 2.1 and Lemma 2.2, it follows that, if (2.1) holds, then

$$A \preceq \text{SR1}(G, A, u) \preceq \text{BFGS}(G, A, u) \preceq \text{DFP}(G, A, u).$$

In other words, the approximation, produced by SR1, is better than the one, produced by BFGS, which is in turn better than the one, produced by DFP.

Let us now justify the efficiency of update (2.2) with $\tau \in [0, 1]$ in ensuring convergence $G \rightarrow A$. For this, we introduce the following measure of progress:

$$\sigma_A(G) \stackrel{\text{def}}{=} \langle A^{-1}, G - A \rangle \stackrel{(1.3)}{=} \langle A^{-1}, G \rangle - n. \quad (2.9)$$

Thus, $\sigma_A(G)$ is the sum of the eigenvalues of the difference $G - A$, measured with respect to the operator A . Clearly, for G , satisfying (2.1), we have $\sigma_A(G) \geq 0$ with $\sigma_A(G) = 0$ if and only if $G = A$. Therefore, we need to ensure that $\sigma_A(G) \rightarrow 0$ by choosing an appropriate update direction u .

First, let us estimate the decrease in the measure σ_A for an arbitrary direction.

Lemma 2.3 *Let (2.1) hold. Then, for any $u \in \mathbb{E}$ and any $\tau \in [0, 1]$, we have*

$$\sigma_A(G) - \sigma_A(\text{Broyd}_\tau(G, A, u)) \geq \frac{\langle (G-A)u, u \rangle}{\langle Au, u \rangle}. \quad (2.10)$$

Proof: Denote $G_+ \stackrel{\text{def}}{=} \text{Broyd}_\tau(G, A, u)$ and assume that $Gu \neq Au$ since otherwise the claim is trivial. By Lemma 2.1, we have

$$G - G_+ \succeq G - \text{DFP}(G, A, u) \stackrel{(2.4)}{=} \frac{Auu^*G + Guu^*A}{\langle Au, u \rangle} - \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \frac{Auu^*A}{\langle Au, u \rangle}.$$

Therefore,

$$\begin{aligned} \sigma_A(G) - \sigma_A(G_+) &\stackrel{(2.9)}{=} \langle A^{-1}, G - G_+ \rangle \geq 2 \frac{\langle Gu, u \rangle}{\langle Au, u \rangle} - \left(\frac{\langle Gu, u \rangle}{\langle Au, u \rangle} + 1 \right) \\ &= \frac{\langle Gu, u \rangle}{\langle Au, u \rangle} - 1 = \frac{\langle (G-A)u, u \rangle}{\langle Au, u \rangle}. \quad \square \end{aligned}$$

According to Lemma 2.3, the choice of the updating direction u directly influences the decrease in the measure σ_A . Ideally, we would like to select a direction u , which maximizes the right-hand side in (2.10). However, this requires finding an eigenvector, corresponding to the maximal eigenvalue of G with respect to A , which might be computationally a difficult problem. Therefore, let us consider another approach.

Let us fix in the space $\mathbb{E}$ some basis:

$$e_1, \dots, e_n \in \mathbb{E}.$$

With respect to this basis, we can define the following greedily selected direction:

$$\bar{u}_A(G) \stackrel{\text{def}}{=} \underset{u \in \{e_1, \dots, e_n\}}{\text{argmax}} \frac{\langle (G-A)u, u \rangle}{\langle Au, u \rangle} = \underset{u \in \{e_1, \dots, e_n\}}{\text{argmax}} \frac{\langle Gu, u \rangle}{\langle Au, u \rangle}. \quad (2.11)$$

Thus, $\bar{u}_A(G)$ is a basis vector, which maximizes the right-hand side in (2.10). Note that for certain choices of the basis, the computation of $\bar{u}_A(G)$ might be relatively simple. For example, if $\mathbb{E} = \mathbb{R}^n$, and $e_1, \dots, e_n$ are coordinate orths, then the calculation of $\bar{u}_A(G)$ requires computing only the *diagonals* of the matrix representations of the operators G and A . The update (2.2), applying the rule (2.11), is called the *greedy quasi-Newton update*.

Let us show that the greedy quasi-Newton update decreases the measure σ_A with a *linear* rate. For this, define

$$B \stackrel{\text{def}}{=} \left(\sum_{i=1}^n e_i e_i^* \right)^{-1}. \quad (2.12)$$

Note that B is a self-adjoint positive definite linear operator from $\mathbb{E}$ to $\mathbb{E}^*$.

Theorem 2.1 *Let (2.1) hold, and let $\mu, L > 0$ be such that*

$$\mu B \preceq A \preceq LB. \quad (2.13)$$

Then, for any $\tau \in [0, 1]$, we have

$$\sigma_A(\text{Broyd}_\tau(G, A, \bar{u}_A(G))) \leq \left(1 - \frac{\mu}{nL}\right) \sigma_A(G). \quad (2.14)$$

Proof: Denote $G_+ \stackrel{\text{def}}{=} \text{Broyd}_\tau(G, A, \bar{u}_A(G))$, and $R \stackrel{\text{def}}{=} G - A$. Applying Lemma 2.3, we obtain

$$\begin{aligned}
\sigma_A(G) - \sigma_A(G_+) &\geq \frac{\langle R\bar{u}_A(G), \bar{u}_A(G) \rangle}{\langle A\bar{u}_A(G), \bar{u}_A(G) \rangle} \stackrel{(2.11)}{=} \max_{1 \leq i \leq n} \frac{\langle Re_i, e_i \rangle}{\langle Ae_i, e_i \rangle} \stackrel{(2.13)}{\geq} \frac{1}{L} \max_{1 \leq i \leq n} \langle Re_i, e_i \rangle \\
&\geq \frac{1}{nL} \sum_{i=1}^n \langle Re_i, e_i \rangle \stackrel{(1.2)}{=} \frac{1}{nL} \sum_{i=1}^n \langle e_i e_i^*, R \rangle \stackrel{(2.12)}{=} \frac{1}{nL} \langle B^{-1}, R \rangle \\
&\stackrel{(2.13)}{\geq} \frac{\mu}{nL} \langle A^{-1}, R \rangle \stackrel{(2.9)}{=} \frac{\mu}{nL} \sigma_A(G). \quad \square
\end{aligned}$$

Remark 2.1 A simple modification of the above proof shows that the factor nL in (2.14) can be improved up to $\langle B^{-1}, A \rangle$. However, to simplify the future analysis, we prefer to work directly with constant L .

3 Unconstrained Quadratic Minimization

Let us demonstrate how we can apply the quasi-Newton updates, described in the previous section, for minimizing the quadratic function

$$f(x) \stackrel{\text{def}}{=} \frac{1}{2} \langle Ax, x \rangle - \langle b, x \rangle, \quad x \in \mathbb{E}, \quad (3.1)$$

where $A : \mathbb{E} \rightarrow \mathbb{E}^*$ is a self-adjoint positive definite linear operator, and $b \in \mathbb{E}^*$.

Let B be the operator, defined in (2.12), and let $\mu, L > 0$ be such that

$$\mu B \preceq A \preceq LB. \quad (3.2)$$

Thus, μ is the *constant of strong convexity* of f , and L is the *Lipschitz constant* of the gradient of f , both measured with respect to the operator B .

Consider the following quasi-Newton scheme:

Initialization: Choose $x_0 \in \mathbb{E}$. Set $G_0 = LB$.	
For $k \geq 0$ iterate: 1. Update $x_{k+1} = x_k - G_k^{-1} \nabla f(x_k)$. 2. Choose $u_k \in \mathbb{E}$ and $\tau_k \in [0, 1]$. 3. Compute $G_{k+1} = \text{Broyd}_{\tau_k}(G_k, A, u_k)$.	(3.3)

Note that scheme (3.3) starts with $G_0 = LB$. Therefore, its first iteration is identical to that one of the standard *gradient method*:

$$x_1 = x_0 - \frac{1}{L} B^{-1} \nabla f(x_0).$$

Also, from (3.2), it follows that $A \preceq G_0$. Hence, in view of Lemma 2.2, we have

$$A \preceq G_k \quad (3.4)$$

for all $k \geq 0$. In particular, all G_k are positive definite, and scheme (3.3) is well-defined.

Remark 3.1 *For avoiding the $O(n^3)$ complexity for computing $G_k^{-1} \nabla f(x_k)$, it is typical for practical implementation of scheme (3.3) to work directly with the inverse operators G_k^{-1} (or, alternatively, with the Cholesky decomposition of G_k). Due to a low-rank structure of the updates (2.2), it is possible to compute efficiently G_{k+1}^{-1} via G_k^{-1} at the cost $O(n^2)$.*

To estimate the convergence rate of scheme (3.3), let us look at the norm of the gradient of f , measured with respect to A :

$$\lambda_f(x) \stackrel{\text{def}}{=} \|\nabla f(x)\|_A^* \stackrel{(1.1)}{=} \langle \nabla f(x), A^{-1} \nabla f(x) \rangle^{1/2}, \quad x \in \mathbb{E}. \quad (3.5)$$

Note that this measure of optimality is directly related to the functional residual. Indeed, let $x^* = A^{-1}b$ be the minimizer of (3.1). Then, using Taylor's formula, we obtain

$$\begin{aligned} f(x) - f^* &= \frac{1}{2} \langle A(x - x^*), x - x^* \rangle = \frac{1}{2} \langle Ax - b, A^{-1}(Ax - b) \rangle \\ &\stackrel{(3.1)}{=} \frac{1}{2} \langle \nabla f(x), A^{-1} \nabla f(x) \rangle \stackrel{(3.5)}{=} \frac{1}{2} \lambda_f^2(x). \end{aligned}$$

The following lemma shows how λ_f changes after one iteration of process (3.3).

Lemma 3.1 *Let $k \geq 0$, and let $\eta_k \geq 1$ be such that*

$$G_k \preceq \eta_k A. \quad (3.6)$$

Then,

$$\lambda_f(x_{k+1}) \leq \left(1 - \frac{1}{\eta_k}\right) \lambda_f(x_k) = \frac{\eta_k - 1}{\eta_k} \lambda_f(x_k).$$

Proof: By Taylor's formula, we have

$$\nabla f(x_{k+1}) = \nabla f(x_k) + A(x_{k+1} - x_k) \stackrel{(3.3)}{=} A(A^{-1} - G_k^{-1}) \nabla f(x_k).$$

Therefore,

$$\lambda_f(x_{k+1}) \stackrel{(3.5)}{=} \langle \nabla f(x_k), (A^{-1} - G_k^{-1}) A (A^{-1} - G_k^{-1}) \nabla f(x_k) \rangle^{1/2}.$$

Note that

$$\frac{1}{\eta_k} A^{-1} \stackrel{(3.6)}{\preceq} G_k^{-1} \stackrel{(3.4)}{\preceq} A^{-1}.$$

Therefore,

$$0 \preceq A^{-1} - G_k^{-1} \preceq \left(1 - \frac{1}{\eta_k}\right) A^{-1}. \quad (3.7)$$

Consequently,

$$(A^{-1} - G_k^{-1})A(A^{-1} - G_k^{-1}) \preceq \left(1 - \frac{1}{\eta_k}\right)^2 A^{-1},$$

and

$$\lambda_f(x_{k+1}) \stackrel{(3.7)}{\leq} \left(1 - \frac{1}{\eta_k}\right) \langle \nabla f(x_k), A^{-1} \nabla f(x_k) \rangle^{1/2} \stackrel{(3.5)}{=} \left(1 - \frac{1}{\eta_k}\right) \lambda_f(x_k). \square$$

Thus, to estimate how fast $\lambda_f(x_k)$ converges to zero, we need to upper bound η_k . There are two ways to proceed, depending on the choice of directions u_k in (3.3).

First, consider the general situation, when we do not impose any restrictions on u_k . In this case, we can guarantee that η_k stays uniformly bounded, and $\lambda_f(x_k) \rightarrow 0$ at a *linear* rate.

Theorem 3.1 *For all $k \geq 0$, in scheme (3.3), we have*

$$A \preceq G_k \preceq \frac{L}{\mu} A, \quad (3.8)$$

and

$$\lambda_f(x_k) \leq \left(1 - \frac{\mu}{L}\right)^k \lambda_f(x_0). \quad (3.9)$$

Proof: Since $G_0 = LB$, in view of (3.2), we have

$$A \preceq G_0 \preceq \frac{L}{\mu} A.$$

By Lemma 2.2, this implies (3.8). Applying now Lemma 3.1, we obtain

$$\lambda_f(x_{k+1}) \leq \left(1 - \frac{\mu}{L}\right) \lambda_f(x_k)$$

for all $k \geq 0$, and (3.9) follows. $\square$

Note that (3.9) is exactly the convergence rate of the standard gradient method. Thus, according to Theorem 3.1, the convergence rate of scheme (3.3) is at least as good as that of the gradient method.

Now assume that the directions u_k in scheme (3.3) are chosen in accordance to the greedy strategy (2.11). Recall that, in this case, we can guarantee that $G_k \rightarrow A$. Therefore, we can expect faster convergence from scheme (3.3).

Theorem 3.2 *Suppose that, for each $k \geq 0$, we choose $u_k = \bar{u}_A(G_k)$ in scheme (3.3). Then, for all $k \geq 0$, we have*

$$A \preceq G_k \preceq \left(1 + \left(1 - \frac{\mu}{nL}\right)^k \frac{nL}{\mu}\right) A, \quad (3.10)$$

and

$$\lambda_f(x_{k+1}) \leq \left(1 - \frac{\mu}{nL}\right)^k \frac{nL}{\mu} \cdot \lambda_f(x_k). \quad (3.11)$$

Proof: We already know that $A \preceq G_k$. Hence,

$$G_k - A \stackrel{(1.4)}{\preceq} \langle A^{-1}, G_k - A \rangle A \stackrel{(2.9)}{=} \sigma_A(G_k)A,$$

or, equivalently,

$$G_k \preceq (1 + \sigma_A(G_k))A.$$

At the same time, by Theorem 2.1, we have

$$\sigma_A(G_k) \leq \left(1 - \frac{\mu}{nL}\right)^k \sigma_A(G_0).$$

Note that

$$\sigma_A(G_0) \stackrel{(2.9)}{=} \langle A^{-1}, G_0 \rangle - n \stackrel{(3.8)}{\leq} \langle A^{-1}, \frac{L}{\mu} A \rangle - n \stackrel{(1.3)}{=} n \left(\frac{L}{\mu} - 1 \right) \leq \frac{nL}{\mu}.$$

Thus, (3.10) is proved. Applying now Lemma 3.1, we obtain (3.11). $\square$

Theorem 3.2 shows that the convergence rate of $\lambda_f(x_k)$ is *superlinear*. Let us now combine this result with Theorem 3.1 and write down the final efficiency estimate. Denote by $k_0 \geq 0$ the number of the first iteration, for which

$$\left(1 - \frac{\mu}{nL}\right)^{k_0} \frac{nL}{\mu} \leq \frac{1}{2}. \quad (3.12)$$

Clearly, $k_0 \leq \frac{nL}{\mu} \ln \frac{2nL}{\mu}$. According to Theorem 2.1, during the first k_0 iterations,

$$\lambda_f(x_k) \leq \left(1 - \frac{\mu}{L}\right)^k \lambda_f(x_0). \quad (3.13)$$

After that, by Theorem 3.2, for all $k \geq 0$, we have

$$\lambda_f(x_{k_0+k+1}) \stackrel{(3.11)}{\leq} \left(1 - \frac{\mu}{nL}\right)^{k_0+k} \frac{nL}{\mu} \lambda_f(x_{k_0+k}) \stackrel{(3.12)}{\leq} \left(1 - \frac{\mu}{nL}\right)^k \frac{1}{2} \lambda_f(x_{k_0+k}),$$

or, more explicitly,

$$\begin{aligned} \lambda_f(x_{k_0+k}) &\leq \lambda_f(x_{k_0}) \prod_{i=0}^{k-1} \left[\left(1 - \frac{\mu}{nL}\right)^i \frac{1}{2} \right] = \left(1 - \frac{\mu}{nL}\right)^{\sum_{i=0}^{k-1} i} \left(\frac{1}{2}\right)^k \lambda_f(x_{k_0}) \\ &= \left(1 - \frac{\mu}{nL}\right)^{\frac{k(k-1)}{2}} \left(\frac{1}{2}\right)^k \lambda_f(x_{k_0}) \\ &\stackrel{(3.13)}{\leq} \left(1 - \frac{\mu}{nL}\right)^{\frac{k(k-1)}{2}} \left(\frac{1}{2}\right)^k \left(1 - \frac{\mu}{L}\right)^{k_0} \lambda_f(x_0). \end{aligned}$$

Note that the first factor in this estimate depends on the *square* of the iteration counter.

To conclude, let us mention one important property of scheme (3.3) with greedily selected u_k . It turns out that, in the particular case, when $\tau_k = 0$ for all $k \geq 0$, i.e. when scheme (3.3) corresponds to the greedy *SR1* method, it will identify the operator A , and consequently, the minimizer x^* of (3.1), in a *finite* number of steps.

Theorem 3.3 *Suppose that, in scheme (3.3), for each $k \geq 0$, we choose $u_k = \bar{u}_A(G_k)$ and $\tau_k = 0$. Then $G_k = A$ for some $0 \leq k \leq n$.*

Proof: Suppose that $R_k \stackrel{\text{def}}{=} G_k - A \neq 0$ for all $0 \leq k \leq n$. Since $R_k \succeq 0$ (see (3.4)), we must have $u_k \notin \text{Ker}(R_k)$ in view of (2.11), and

$$R_{k+1} \stackrel{(2.3)}{=} R_k - \frac{R_k u_k u_k^* R_k}{\langle R_k u_k, u_k \rangle}$$

for all $0 \leq k \leq n$. From this formula, it is easily seen that

- (1) $\text{Ker}(R_k) \subseteq \text{Ker}(R_{k+1})$,
- (2) $u_k \in \text{Ker}(R_{k+1})$.

Thus, the dimension of $\text{Ker}(R_k)$ grows at least by 1 at every iteration. In particular, the dimension of $\text{Ker}(R_{n+1})$ must be at least $n+1$, which is impossible, since the operator R_{n+1} acts in an n -dimensional vector space. $\square$

Note that for other updates, such as (2.4) or (2.6), the inclusion $\text{Ker}(R_k) \subseteq \text{Ker}(R_{k+1})$ is, in general, no longer valid.

4 Minimization of General Functions

Now consider a general unconstrained minimization problem:

$$\min_{x \in \mathbb{E}} f(x), \tag{4.1}$$

where $f : \mathbb{E} \rightarrow \mathbb{R}$ is a twice differentiable function with positive definite Hessian. Our goal is to extend the results, obtained in the previous section, onto the problem (4.1), assuming that the methods can start from a sufficiently good initial point x_0 .

Our main assumption is that the Hessians of f are close to each other in the sense that there exists a constant $M \geq 0$, such that

$$\nabla^2 f(y) - \nabla^2 f(x) \preceq M \|y - x\|_z \nabla^2 f(w) \tag{4.2}$$

for all $x, y, z, w \in \mathbb{E}$. We call such a function f *strongly self-concordant*. Note that strongly self-concordant functions form a subclass of self-concordant functions. Indeed, let us choose a point $x \in \mathbb{E}$ and a direction $h \in \mathbb{E}$. Then, for all $t > 0$, we have

$$\langle [\nabla^2 f(x + th) - \nabla^2 f(x)]h, h \rangle \leq Mt \|h\|_x^3.$$

Dividing this inequality by t and computing the limit as $t \downarrow 0$, we obtain

$$D^3 f(x)[h, h, h] \leq M \|h\|_x^3.$$

for all $h \in \mathbb{E}$. Thus, function f is self-concordant with constant $\frac{1}{2}M$ (see [19, 27]).

The main example of a strongly self-concordant function is a strongly convex function with Lipschitz continuous Hessian. Note however that strong self-concordancy is an *affine-invariant* property.

Example 4.1 Let $C : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint positive definite operator. Suppose there exist $\beta > 0$ and $L_2 \geq 0$, such that the function f is β -strongly convex and its Hessian is L_2 -Lipschitz continuous with respect to the norm $\|\cdot\|_C$. Then f is strongly self-concordant with constant $M = \frac{L_2}{\beta^{3/2}}$.

Proof: By strong convexity of f , we have

$$\beta C \preceq \nabla^2 f(x) \quad (4.3)$$

for all $x \in \mathbb{E}$. Therefore, using the Lipschitz continuity of the Hessian, we obtain

$$\begin{aligned} \nabla^2 f(y) - \nabla^2 f(x) &\preceq L_2 \|y - x\|_C C \stackrel{(1.1)}{=} L_2 \langle C(y - x), y - x \rangle^{1/2} C \\ &\stackrel{(4.3)}{\preceq} \frac{L_2}{\mu^{1/2}} \langle \nabla^2 f(z)(y - x), y - x \rangle^{1/2} C \\ &\stackrel{(1.1)}{=} \frac{L_2}{\mu^{1/2}} \|y - x\|_z C \stackrel{(4.3)}{\preceq} \frac{L_2}{\mu^{3/2}} \|y - x\|_z \nabla^2 f(w) \end{aligned}$$

for all $x, y, z, w \in \mathbb{E}$. □

Let us establish some useful relations for strongly self-concordant functions.

Lemma 4.1 Let $x, y \in \mathbb{E}$, and let $r \stackrel{\text{def}}{=} \|y - x\|_x$. Then,

$$\frac{\nabla^2 f(x)}{1+Mr} \preceq \nabla^2 f(y) \preceq (1+Mr) \nabla^2 f(x). \quad (4.4)$$

Also, for $J \stackrel{\text{def}}{=} \int_0^1 \nabla^2 f(x + t(y - x)) dt$, we have

$$\frac{\nabla^2 f(x)}{1+\frac{Mr}{2}} \preceq J \preceq \left(1 + \frac{Mr}{2}\right) \nabla^2 f(x), \quad (4.5)$$

and

$$\frac{\nabla^2 f(y)}{1+\frac{Mr}{2}} \preceq J \preceq \left(1 + \frac{Mr}{2}\right) \nabla^2 f(y). \quad (4.6)$$

Proof: Denote $h \stackrel{\text{def}}{=} y - x$. Taking $z = w = x$ in (4.2), we obtain

$$\nabla^2 f(y) - \nabla^2 f(x) \preceq Mr \nabla^2 f(x),$$

which gives us the second relation in (4.4) after moving $\nabla^2 f(x)$ into the right-hand side. Interchanging now x and y in (4.2) and taking $z = x$, $w = y$, we get

$$\nabla^2 f(x) - \nabla^2 f(y) \preceq Mr \nabla^2 f(y),$$

which gives us the first relation in (4.4) after moving $\nabla^2 f(x)$ into the right-hand side and then dividing by $1 + Mr$.

Choosing now $y = x + th$ in (4.2) for $t > 0$, and $w = z = x$, we obtain

$$\nabla^2 f(x + th) - \nabla^2 f(x) \preceq M \|th\|_x \nabla^2 f(x) = Mrt \nabla^2 f(x).$$

This gives us the second relation in (4.5) after integrating for t from 0 to 1 and moving $\nabla^2 f(x)$ into the right-hand side. Interchanging x and y in (4.2) and taking $y = x + th$ for $t > 0$, $z = x$, while leaving w arbitrary, we get

$$\nabla^2 f(x) - \nabla^2 f(x + th) \preceq M\| - th\|_x \nabla^2 f(w) = Mrt \nabla^2 f(w).$$

Hence, by integrating for t from 0 to 1,

$$\nabla^2 f(x) - J \preceq \frac{Mr}{2} \nabla^2 f(w).$$

Taking now $w = x + th$ and integrating again, we obtain

$$\nabla^2 f(x) - J \preceq \frac{Mr}{2} \int_0^1 \nabla^2 f(x + th) dt = \frac{Mr}{2} J,$$

and the first inequality in (4.5) follows after moving J to the right-hand side and dividing by $1 + \frac{Mr}{2}$.

Relations (4.6) can be proved similarly to (4.5). $\square$

Let us now estimate the progress of a general quasi-Newton step. As before, for measuring the progress, we use the *local norm of the gradient*:

$$\lambda_f(x) \stackrel{\text{def}}{=} \|\nabla f(x)\|_x^* \stackrel{(1.1)}{=} \langle \nabla f(x), \nabla^2 f(x)^{-1} \nabla f(x) \rangle^{1/2}, \quad x \in \mathbb{E}. \quad (4.7)$$

Lemma 4.2 *Let $x \in \mathbb{E}$, and let $G : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint linear operator, such that*

$$\nabla^2 f(x) \preceq G \preceq \eta \nabla^2 f(x) \quad (4.8)$$

for some $\eta \geq 1$. Let

$$x_+ \stackrel{\text{def}}{=} x - G^{-1} \nabla f(x), \quad (4.9)$$

and let $\lambda \stackrel{\text{def}}{=} \lambda_f(x)$ be such that $M\lambda \leq 2$. Then, $r \stackrel{\text{def}}{=} \|x_+ - x\|_x \leq \lambda$, and

$$\lambda_f(x_+) \leq \left(1 + \frac{M\lambda}{2}\right) \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta} \lambda.$$

Proof: Denote $J \stackrel{\text{def}}{=} \int_0^1 \nabla^2 f(x + t(x_+ - x)) dt$. Applying Taylor's formula, we obtain

$$\nabla f(x_+) = \nabla f(x) + J(x_+ - x) \stackrel{(4.9)}{=} J(J^{-1} - G^{-1}) \nabla f(x). \quad (4.10)$$

Note that

$$\begin{aligned} r &= \|x_+ - x\|_x \stackrel{(4.9)}{=} \|G^{-1} \nabla f(x)\|_x \stackrel{(1.1)}{=} \langle \nabla f(x), G^{-1} \nabla^2 f(x) G^{-1} \nabla f(x) \rangle^{1/2} \\ &\stackrel{(4.8)}{\leq} \langle \nabla f(x), G^{-1} \nabla f(x) \rangle^{1/2} \stackrel{(4.8)}{\leq} \langle \nabla f(x), \nabla^2 f(x)^{-1} \nabla f(x) \rangle^{1/2} \stackrel{(4.7)}{=} \lambda. \end{aligned}$$

Hence, in view of Lemma 4.1, we have

$$\frac{\nabla^2 f(x)}{1 + \frac{M\lambda}{2}} \preceq J \preceq \left(1 + \frac{M\lambda}{2}\right) \nabla^2 f(x), \quad J \preceq \left(1 + \frac{M\lambda}{2}\right) \nabla^2 f(x_+). \quad (4.11)$$

Therefore,

$$\begin{aligned}
\lambda_f(x_+) &\stackrel{(4.7)}{=} \langle \nabla f(x_+), \nabla^2 f(x_+)^{-1} \nabla f(x_+) \rangle^{1/2} \\
&\stackrel{(4.11)}{\leq} \sqrt{1 + \frac{M\lambda}{2}} \langle \nabla f(x_+), J^{-1} \nabla f(x_+) \rangle^{1/2} \\
&\stackrel{(4.10)}{=} \sqrt{1 + \frac{M\lambda}{2}} \langle \nabla f(x), (J^{-1} - G^{-1})J(J^{-1} - G^{-1}) \nabla f(x) \rangle^{1/2}.
\end{aligned} \tag{4.12}$$

Further,

$$\frac{1}{1 + \frac{M\lambda}{2}} J \stackrel{(4.11)}{\preceq} \nabla^2 f(x) \stackrel{(4.8)}{\preceq} G \stackrel{(4.8)}{\preceq} \eta \nabla^2 f(x) \stackrel{(4.11)}{\preceq} \eta \left(1 + \frac{M\lambda}{2}\right) J.$$

Hence,

$$\frac{1}{(1 + \frac{M\lambda}{2})\eta} J^{-1} \preceq G^{-1} \preceq \left(1 + \frac{M\lambda}{2}\right) J^{-1},$$

and

$$-\left(1 - \frac{1}{(1 + \frac{M\lambda}{2})\eta}\right) J^{-1} \preceq G^{-1} - J^{-1} \preceq \frac{M\lambda}{2} J^{-1}.$$

Note that

$$1 - \frac{1}{(1 + \frac{M\lambda}{2})\eta} \leq 1 - \frac{1 - \frac{M\lambda}{2}}{\eta} = \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta},$$

and, since $M\lambda \leq 2$,

$$\frac{M\lambda}{2} = 1 - \left(1 - \frac{M\lambda}{2}\right) \leq 1 - \frac{1 - \frac{M\lambda}{2}}{\eta} = \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta}.$$

Therefore,

$$-\frac{\eta - 1 + \frac{M\lambda}{2}}{\eta} J^{-1} \preceq G^{-1} - J^{-1} \preceq \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta} J^{-1}.$$

Consequently,

$$(G^{-1} - J^{-1})J(G^{-1} - J^{-1}) \preceq \left(\frac{\eta - 1 + \frac{M\lambda}{2}}{\eta}\right)^2 J^{-1}.$$

and thus,

$$\begin{aligned}
\lambda_f(x_+) &\stackrel{(4.12)}{\leq} \sqrt{1 + \frac{M\lambda}{2} \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta}} \langle \nabla f(x), J^{-1} \nabla f(x) \rangle^{1/2} \\
&\stackrel{(4.11)}{\leq} \left(1 + \frac{M\lambda}{2}\right) \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta} \langle \nabla f(x), \nabla^2 f(x)^{-1} \nabla f(x) \rangle^{1/2} \\
&\stackrel{(4.7)}{=} \left(1 + \frac{M\lambda}{2}\right) \frac{\eta - 1 + \frac{M\lambda}{2}}{\eta} \lambda. \quad \square
\end{aligned}$$

Now we need to analyze what happens with the Hessian approximation after a quasi-Newton update. Let G be the current approximation of $\nabla^2 f(x)$, satisfying, as usual, the condition

$$\nabla^2 f(x) \preceq G. \tag{4.13}$$

Using this approximation, we can compute the new test point

$$x_+ = x - G^{-1} \nabla f(x).$$

After that, we would like to update G into a new operator G_+ , approximating the Hessian $\nabla^2 f(x_+)$ at the new point and satisfying the condition

$$\nabla^2 f(x_+) \preceq G_+.$$

A natural idea is, of course, to set

$$G_+ = \text{Broyd}_\tau(G, \nabla^2 f(x_+), u) \quad (4.14)$$

for some $u \in \mathbb{E}$ and $\tau \in [0, 1]$. However, we cannot do this, since update (4.14) is well-defined only when

$$\nabla^2 f(x_+) \preceq G$$

(see Section 2), which may not be true, even though (4.13) holds. To avoid this problem, let us apply the following *correction strategy*:

1. Choose some $\delta \geq 0$, and set $\tilde{G} = (1 + \delta)G$.
2. Compute G_+ , using (4.14) with G replaced by $\tilde{G}$.

Clearly, for some value of δ , the condition $\nabla^2 f(x_+) \preceq \tilde{G}$ will be valid. If, at the same time, this δ is sufficiently small, then the above correction strategy should not introduce too big error.

Lemma 4.3 *Let $x \in \mathbb{E}$, and let $G : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint linear operator, such that*

$$\nabla^2 f(x) \preceq G \preceq \eta \nabla^2 f(x) \quad (4.15)$$

for some $\eta \geq 1$. Let $x_+ \in \mathbb{E}$, let $r \stackrel{\text{def}}{=} \|x_+ - x\|_x$. Then

$$\tilde{G} \stackrel{\text{def}}{=} (1 + Mr)G \succeq \nabla^2 f(x_+), \quad (4.16)$$

and, for all $u \in \mathbb{E}$ and $\tau \in [0, 1]$, we have

$$\nabla^2 f(x_+) \preceq \text{Broyd}_\tau(\tilde{G}, \nabla^2 f(x_+), u) \preceq [(1 + Mr)^2 \eta] \nabla^2 f(x_+),$$

Proof: Note that

$$\nabla^2 f(x_+) \stackrel{(4.4)}{\preceq} (1 + Mr) \nabla^2 f(x) \stackrel{(4.15)}{\preceq} (1 + Mr)G = \tilde{G},$$

and,

$$\tilde{G} = (1 + Mr)G \stackrel{(4.15)}{\preceq} (1 + Mr)\eta \nabla^2 f(x) \stackrel{(4.4)}{\preceq} (1 + Mr)^2 \eta \nabla^2 f(x_+).$$

Thus,

$$\nabla^2 f(x_+) \preceq \tilde{G} \preceq (1 + Mr)^2 \eta \nabla^2 f(x_+),$$

and the claim now follows from Lemma 2.2. $\square$

Let us now make one more assumption about the function f . We assume that, with respect to the operator B , defined by (2.12), the function f is *strongly convex*, and its gradient is *Lipschitz continuous*, i.e. there exist $\mu, L > 0$, such that, for all $x, y \in \mathbb{E}$, we have

$$\mu B \preceq \nabla^2 f(x) \preceq LB. \quad (4.17)$$

Remark 4.1 *In fact, for our purposes, it is enough to require that conditions (4.2), (4.17) hold only in a neighborhood of a solution, but, for the sake of simplicity, we do not do this.*

We are ready to write down the scheme of our quasi-Newton methods. For simplicity, we assume that the constants M and L are available.

Initialization: Choose $x_0 \in \mathbb{E}$. Set $G_0 = LB$.

For $k \geq 0$ iterate:

1. Update $x_{k+1} = x_k - G_k^{-1} \nabla f(x_k)$.
 2. Compute $r_k = \|x_{k+1} - x_k\|_{x_k}$ and set $\tilde{G}_k = (1 + Mr_k)G_k$.
 2. Choose $u_k \in \mathbb{E}$ and $\tau_k \in [0, 1]$.
 4. Compute $G_{k+1} = \text{Broyd}_{\tau_k}(\tilde{G}_k, \nabla^2 f(x_{k+1}), u_k)$.
- (4.18)

Remark 4.2 *Similarly to Remark 3.1, in a practical implementation of scheme (4.18), one should work directly with the inverse operators G_k^{-1} , or with the Cholesky decomposition of G_k . Note that the correction step $\tilde{G}_k = (1 + Mr_k)G_k$ does not affect the complexity of the iteration.*

As before, we present two convergence results for scheme (4.18). The first one establishes linear convergence and can be seen as a generalization of Theorem 3.1. Note that for this result the directions u_k in the method (4.18) can be chosen arbitrarily.

Theorem 4.1 *Suppose the initial point x_0 is chosen sufficiently close to the solution:*

$$M\lambda_f(x_0) \leq \frac{\ln \frac{3}{2}}{4} \frac{\mu}{L}. \quad (4.19)$$

Then, for all $k \geq 0$, we have

$$\nabla^2 f(x_k) \preceq G_k \preceq e^{2M \sum_{i=0}^{k-1} \lambda_f(x_i) \frac{L}{\mu}} \nabla^2 f(x_k) \preceq \frac{3L}{2\mu} \nabla^2 f(x_k), \quad (4.20)$$

and

$$\lambda_f(x_k) \leq \left(1 - \frac{\mu}{2L}\right)^k \lambda_f(x_0). \quad (4.21)$$

Proof: In view of (4.17), we have

$$\nabla^2 f(x_0) \preceq G_0 \preceq \frac{L}{\mu} \nabla^2 f(x_0).$$

Therefore, for $k = 0$, both (4.20) and (4.21) are satisfied.

Now let $k \geq 0$, and suppose (4.20), (4.21) have already been proved for all $0 \leq k' \leq k$. Denote $\lambda_k \stackrel{\text{def}}{=} \lambda_f(x_k)$, $r_k \stackrel{\text{def}}{=} \|x_{k+1} - x_k\|_{x_k}$, and

$$\eta_k \stackrel{\text{def}}{=} e^{2M \sum_{i=0}^{k-1} \lambda_i \frac{L}{\mu}}. \quad (4.22)$$

Note that

$$M \sum_{i=0}^k \lambda_i \stackrel{(4.21)}{\leq} M \lambda_0 \sum_{i=0}^k \left(1 - \frac{\mu}{2L}\right)^i \leq \frac{2L}{\mu} M \lambda_0 \stackrel{(4.19)}{\leq} \frac{\ln \frac{3}{2}}{2}. \quad (4.23)$$

Applying Lemma 4.2, we obtain that

$$r_k \leq \lambda_k \quad (4.24)$$

and

$$\lambda_{k+1} \leq \left(1 + \frac{M\lambda_k}{2}\right) \frac{\eta_k^{-1} + \frac{M\lambda_k}{2}}{\eta_k} \lambda_k = \left(1 + \frac{M\lambda_k}{2}\right) \left(1 - \frac{1 - \frac{M\lambda_k}{2}}{\eta_k}\right) \lambda_k. \quad (4.25)$$

Note (using the inequality $1 - t \geq e^{-2t}$, valid at least for $0 \leq t \leq \frac{1}{2}$) that

$$\frac{1 - \frac{M\lambda_k}{2}}{\eta_k} \geq e^{-M\lambda_k \eta_k^{-1}} \stackrel{(4.22)}{=} e^{-M\lambda_k - 2M \sum_{i=0}^{k-1} \lambda_i \frac{L}{\mu}} \geq e^{-2M \sum_{i=0}^k \lambda_i \frac{L}{\mu}} \stackrel{(4.23)}{\geq} \frac{2\mu}{3L},$$

and also (using $\ln(1+t) \leq t$, valid for $t \geq 0$) that

$$\frac{M\lambda_k}{2} \stackrel{(4.19)}{\leq} \frac{\ln \frac{3}{2} \frac{\mu}{L}}{8} \leq \frac{\mu}{16L}.$$

Hence,

$$\left(1 + \frac{M\lambda_k}{2}\right) \left(1 - \frac{1 - \frac{M\lambda_k}{2}}{\eta_k}\right) \leq \left(1 + \frac{\mu}{16L}\right) \left(1 - \frac{2\mu}{3L}\right) \leq 1 - \left(\frac{2}{3} - \frac{1}{16}\right) \frac{\mu}{L} \leq 1 - \frac{\mu}{2L}.$$

Consequently,

$$\lambda_{k+1} \stackrel{(4.25)}{\leq} \left(1 - \frac{\mu}{2L}\right) \lambda_k \stackrel{(4.21)}{\leq} \left(1 - \frac{\mu}{2L}\right)^{k+1} \lambda_0.$$

Finally, from Lemma 4.3, it follows that

$$\begin{aligned} \nabla^2 f(x_{k+1}) &\preceq G_{k+1} \preceq (1 + Mr_k)^2 \eta_k \nabla^2 f(x_{k+1}) \\ &\stackrel{(4.24)}{\preceq} (1 + M\lambda_k)^2 \eta_k \nabla^2 f(x_{k+1}) \preceq e^{2M\lambda_k \eta_k} \nabla^2 f(x_{k+1}) \\ &\stackrel{(4.22)}{=} e^{2M \sum_{i=0}^k \lambda_i \frac{L}{\mu}} \nabla^2 f(x_{k+1}) \stackrel{(4.23)}{\preceq} \frac{3L}{2\mu} \nabla^2 f(x_{k+1}). \end{aligned} \quad (4.26)$$

Thus, (4.20), (4.21) are valid for $k' = k + 1$, and we can continue by induction. $\square$

Now let us analyze the greedy strategy. First, we analyze how the Hessian approximation measure (2.9) changes after one iteration.

Lemma 4.4 *Let $x \in \mathbb{E}$, and let $G : \mathbb{E} \rightarrow \mathbb{E}^*$ be a self-adjoint linear operator, such that be such that $\nabla^2 f(x) \preceq G$. Let $x_+ \in \mathbb{E}$, let $r \stackrel{\text{def}}{=} \|x_+ - x\|_x$, and let*

$$\tilde{G} \stackrel{\text{def}}{=} (1 + Mr)G. \quad (4.27)$$

Then, for any $\tau \in [0, 1]$, we have

$$\sigma_{x_+}(\text{Broyd}_\tau(\tilde{G}, \nabla^2 f(x_+), \bar{u}_{x_+}(G))) \leq (1 - \frac{\mu}{nL})(1 + Mr)^2 \left(\sigma_x(G) + \frac{2nMr}{1+Mr} \right).$$

Proof: We already know from Lemma 4.3 that $\nabla^2 f(x_+) \preceq \tilde{G}$. Also note that $\bar{u}_{x_+}(\tilde{G}) = \bar{u}_{x_+}(G)$ (see (2.11)). Hence, by Theorem 2.1, we have

$$\sigma_{x_+}(\text{Broyd}_\tau(\tilde{G}, \nabla^2 f(x_+), \bar{u}_{x_+}(G))) \leq (1 - \frac{\mu}{nL})\sigma_{x_+}(\tilde{G}).$$

Further,

$$\begin{aligned} \sigma_{x_+}(\tilde{G}) &\stackrel{(2.9)}{=} \langle \nabla^2 f(x_+)^{-1}, \tilde{G} \rangle - n \stackrel{(4.27)}{=} (1 + Mr) \langle \nabla^2 f(x_+)^{-1}, G \rangle - n \\ &\stackrel{(4.4)}{\leq} (1 + Mr)^2 \langle \nabla^2 f(x)^{-1}, G \rangle - n \stackrel{(2.9)}{=} (1 + Mr)^2 (\sigma_x(G) + n) - n \\ &= (1 + Mr)^2 \sigma_x(G) + n((1 + Mr)^2 - 1) \\ &= (1 + Mr)^2 \sigma_x(G) + 2nMr \left(1 + \frac{Mr}{2}\right) \\ &\leq (1 + Mr)^2 \left(\sigma_x(G) + \frac{2nMr}{1+Mr} \right). \quad \square \end{aligned}$$

Now we can prove superlinear convergence. In what follows, we assume that $n \geq 2$.

Theorem 4.2 *Suppose that, in scheme (4.18), for each $k \geq 0$ we take $u_k = \bar{u}_{x_{k+1}}(G_k)$. And suppose that the initial point x_0 is chosen sufficiently close to the solution:*

$$M\lambda_f(x_0) \leq \frac{\ln 2}{4(2n+1)} \frac{\mu}{L} \quad \left(\leq \frac{\ln \frac{3}{2}}{4} \frac{\mu}{L} \right). \quad (4.28)$$

Then, for all $k \geq 0$, we have

$$\nabla^2 f(x_k) \preceq G_k \preceq \left(1 + \left(1 - \frac{\mu}{nL} \right)^k \frac{2nL}{\mu} \right) \nabla^2 f(x_k), \quad (4.29)$$

and

$$\lambda_f(x_{k+1}) \leq \left(1 - \frac{\mu}{nL} \right)^k \frac{2nL}{\mu} \cdot \lambda_f(x_k). \quad (4.30)$$

Proof: Denote $\lambda_k \stackrel{\text{def}}{=} \lambda_f(x_k)$ and $\sigma_k \stackrel{\text{def}}{=} \sigma_{x_k}(G_k)$ for $k \geq 0$. In view of Theorem 4.1, the first relation in (4.29) is indeed true, and also

$$M \sum_{i=0}^k \lambda_i \leq M \lambda_0 \sum_{i=0}^k \left(1 - \frac{\mu}{2L} \right)^i \leq \frac{2L}{\mu} \lambda_0 \stackrel{(4.28)}{\leq} \frac{\ln 2}{2(2n+1)}. \quad (4.31)$$

for all $k \geq 0$.

Let us show by induction that, for all $k \geq 0$, we have

$$\sigma_k + 2nM\lambda_k \leq \theta_k. \quad (4.32)$$

where

$$\theta_k \stackrel{\text{def}}{=} \left(1 - \frac{\mu}{nL}\right)^k e^{2(2n+1)M \sum_{i=0}^{k-1} \lambda_i \frac{nL}{\mu}} \stackrel{(4.31)}{\leq} \left(1 - \frac{\mu}{nL}\right)^k \frac{2nL}{\mu}. \quad (4.33)$$

Indeed, since $\nabla^2 f(x_0) \preceq G_0 \preceq \frac{L}{\mu} \nabla^2 f(x_0)$ (see (4.17)), we have

$$\begin{aligned} \sigma_0 + 2nM\lambda_0 &\stackrel{(2.9)}{=} \langle \nabla^2 f(x_0)^{-1}, G_0 \rangle - n + 2nM\lambda_0 \\ &\leq \langle \nabla^2 f(x_0)^{-1}, \frac{L}{\mu} \nabla^2 f(x_0) \rangle - n + 2nM\lambda_0 \\ &\stackrel{(1.3)}{=} n \left(\frac{L}{\mu} - 1 \right) + 2nM\lambda_0 \stackrel{(4.28)}{\leq} n \left(\frac{L}{\mu} - 1 \right) + \frac{n \ln 2}{2(2n+1)} \leq \frac{nL}{\mu}. \end{aligned}$$

Therefore, for $k = 0$, inequality (4.32) is satisfied. Now suppose that it is also satisfied for some $k \geq 0$. Since $\nabla^2 f(x_k) \preceq G_k$, we know that

$$G_k - \nabla^2 f(x_k) \stackrel{(1.4)}{\preceq} \sigma_k \nabla^2 f(x_k),$$

or, equivalently,

$$G_k \preceq (1 + \sigma_k) \nabla^2 f(x_k). \quad (4.34)$$

Therefore, applying Lemma 4.2, we obtain that

$$r_k \stackrel{\text{def}}{=} \|x_{k+1} - x_k\|_{x_k} \leq \lambda_k, \quad (4.35)$$

and

$$\begin{aligned} \lambda_{k+1} &\leq \left(1 + \frac{M\lambda_k}{2}\right) \frac{\sigma_k + \frac{M\lambda_k}{2}}{1 + \sigma_k} \lambda_k \leq \left(1 + \frac{M\lambda_k}{2}\right) (\sigma_k + 2nM\lambda_k) \lambda_k \\ &\stackrel{(4.32)}{\leq} \left(1 + \frac{M\lambda_k}{2}\right) \theta_k \lambda_k \leq e^{\frac{M\lambda_k}{2}} \theta_k \lambda_k \leq e^{2M\lambda_k} \theta_k \lambda_k. \end{aligned} \quad (4.36)$$

Further, by Lemma 4.4, we have

$$\begin{aligned} \sigma_{k+1} &\leq \left(1 - \frac{\mu}{nL}\right) (1 + Mr_k)^2 \left(\sigma_k + \frac{2nMr_k}{1 + Mr_k}\right) \\ &\stackrel{(4.35)}{\leq} \left(1 - \frac{\mu}{nL}\right) (1 + M\lambda_k)^2 \left(\sigma_k + \frac{2nM\lambda_k}{1 + M\lambda_k}\right) \\ &\leq \left(1 - \frac{\mu}{nL}\right) (1 + M\lambda_k)^2 (\sigma_k + 2nM\lambda_k) \\ &\stackrel{(4.32)}{\leq} \left(1 - \frac{\mu}{nL}\right) (1 + M\lambda_k)^2 \theta_k \leq \left(1 - \frac{\mu}{nL}\right) e^{2M\lambda_k} \theta_k. \end{aligned}$$

Note that $\frac{1}{2} \leq 1 - \frac{\mu}{nL}$ since $n \geq 2$. Therefore,

$$\begin{aligned} \sigma_{k+1} + 2nM\lambda_{k+1} &\leq \left(1 - \frac{\mu}{nL}\right) e^{2M\lambda_k} \theta_k + e^{2M\lambda_k} \theta_k 2nM\lambda_k \\ &\leq \left(1 - \frac{\mu}{nL}\right) e^{2M\lambda_k} \theta_k + \left(1 - \frac{\mu}{nL}\right) e^{2M\lambda_k} \theta_k 4nM\lambda_k \\ &= \left(1 - \frac{\mu}{nL}\right) e^{2M\lambda_k} (1 + 4nM\lambda_k) \theta_k \\ &\leq \left(1 - \frac{\mu}{nL}\right) e^{2(2n+1)M\lambda_k} \theta_k \stackrel{(4.33)}{=} \theta_{k+1}. \end{aligned}$$

Thus, (4.32) is proved.

Let us fix now some $k \geq 0$. Since $\lambda_k \geq 0$, we have

$$\sigma_k \leq \sigma_k + 2M\lambda_k \stackrel{(4.32)}{\leq} \theta_k \stackrel{(4.33)}{\leq} \left(1 - \frac{\mu}{nL}\right)^k \frac{2nL}{\mu}.$$

This proves the second relation in (4.29) in view of (4.34). Finally,

$$\lambda_{k+1} \stackrel{(4.36)}{\leq} e^{2M\lambda_k} \theta_k \lambda_k \leq e^{2(2n+1)M\lambda_k} \theta_k \lambda_k \stackrel{(4.33)}{=} \frac{\theta_{k+1}}{1 - \frac{\mu}{nL}} \lambda_k \stackrel{(4.33)}{\leq} \left(1 - \frac{\mu}{nL}\right)^k \frac{2nL}{\mu} \lambda_k,$$

and we obtain (4.30). $\square$

Similarly to the quadratic case, combining Theorem 4.1 with Theorem 4.2, we obtain the following final efficiency estimate:

$$\lambda_f(x_{k_0+k}) \leq \left(1 - \frac{\mu}{nL}\right)^{\frac{k(k-1)}{2}} \left(\frac{1}{2}\right)^k \left(1 - \frac{\mu}{2L}\right)^{k_0} \lambda_f(x_0), \quad k \geq 0,$$

where $k_0 \leq \frac{nL}{\mu} \ln \frac{2nL}{\mu}$.

5 Numerical Experiments

In this section, we present preliminary computational results for greedy quasi-Newton methods, applied to the following test function²:

$$f(x) \stackrel{\text{def}}{=} \ln \left(\sum_{j=1}^m e^{\langle c_j, x \rangle - b_j} \right) + \frac{1}{2} \sum_{j=1}^m \langle c_j, x \rangle^2 + \frac{\gamma}{2} \|x\|^2, \quad x \in \mathbb{R}^n, \quad (5.1)$$

where $c_1, \dots, c_m \in \mathbb{R}^n$, $b_1, \dots, b_m \in \mathbb{R}$, and $\gamma > 0$.

We compare scheme (4.18) (which realizes GrDFP, GrBFGS and GrSR1, depending on the choice of τ_k) with the usual gradient method (GM) and standard quasi-Newton methods DFP, BFGS and SR1.

All the standard methods need access only to the gradient of function f :

$$\nabla f(x) = g(x) + \sum_{j=1}^m \langle c_j, x \rangle c_j + \gamma x, \quad g(x) \stackrel{\text{def}}{=} \sum_{j=1}^m \pi_j(x) c_j, \quad (5.2)$$

where

$$\pi_j(x) \stackrel{\text{def}}{=} \frac{e^{\langle c_j, x \rangle - b_j}}{\sum_{j'=1}^m e^{\langle c_{j'}, x \rangle - b_{j'}}} \in [0, 1], \quad j = 1, \dots, m.$$

Note that, for a given point $x \in \mathbb{R}^n$, $\nabla f(x)$ can be computed in $O(mn)$ operations.

²Note that we work in the space $\mathbb{E} = \mathbb{R}^n$ and identify $\mathbb{E}^*$ with $\mathbb{E}$ in such a way that $\langle \cdot, \cdot \rangle$ is the standard dot product, and $\|\cdot\|$ is the standard Euclidean norm. Linear operators from $\mathbb{E}$ to $\mathbb{E}^*$ are identified with $n \times n$ matrices.

For greedy methods, to implement the Hessian approximation update, at every iteration, we need to carry out some additional operations with the Hessian

$$\begin{aligned}\nabla^2 f(x) &= \sum_{j=1}^m \pi_j(x) c_j c_j^T - g(x)g(x)^T + \sum_{j=1}^m c_j c_j^T + \gamma I \\ &= \sum_{j=1}^m (\pi_j(x) + 1) c_j c_j^T - g(x)g(x)^T + \gamma I.\end{aligned}\tag{5.3}$$

Namely, given a point $x \in \mathbb{R}^n$, we need to be able to perform the following two actions:

- For all $1 \leq i \leq n$, compute the values

$$\langle \nabla^2 f(x) e_i, e_i \rangle \stackrel{(5.3)}{=} \sum_{j=1}^m (\pi_j(x) + 1) \langle c_j, e_i \rangle^2 - \langle g(x), e_i \rangle^2 + \gamma,$$

where $e_1, \dots, e_n$ are the basis vectors.

- For a given direction $h \in \mathbb{R}^n$, compute the Hessian-vector product

$$\nabla^2 f(x) h \stackrel{(5.3)}{=} \sum_{j=1}^m (\pi_j(x) + 1) \langle c_j, h \rangle c_j - \langle g(x), h \rangle g(x) + \gamma h.$$

Let us take the basis $e_1, \dots, e_n$, comprised of the standard coordinate directions:

$$e_i \stackrel{\text{def}}{=} (0, \dots, 0, 1, 0, \dots, 0)^T, \quad 1 \leq i \leq n.\tag{5.4}$$

Then, both the above operations have a cost of $O(mn)$. Thus, the cost of one iteration for all the methods under our consideration is comparable.

Note that for basis (5.4), the matrix B , defined by (2.12), is the identity matrix:

$$B = I.$$

Hence, the Lipschitz constant of the gradient of f with respect to B can be taken as follows (see (5.3)):

$$L = 2 \sum_{j=1}^m \|c_j\|^2 + \gamma.$$

All quasi-Newton methods in our comparison start from the same initial Hessian approximation $G_0 = LB$, and use unit step sizes.

Finally, for greedy quasi-Newton methods, we also need to provide an estimate of the strong self-concordancy parameter. Note that, with respect to the operator $\sum_{j=1}^m c_j c_j^T$, the function f is 1-strongly convex and its Hessian is 2-Lipschitz continuous (see e.g. [28, Ex. 1]). Hence, in view of Example 4.1, the strong self-concordancy parameter can be chosen as follows:

$$M = 2.$$

The data, defining the test function (5.1), is randomly generated in the following way. First, we generate a collection of random vectors

$$\hat{c}_1, \dots, \hat{c}_m$$

with entries, uniformly distributed in the interval $[-1, 1]$. Then we generate $b_1, \dots, b_m$ from the same distribution. Using this data, we form a preliminary function

$$\hat{f}(x) \stackrel{\text{def}}{=} \ln \left(\sum_{j=1}^m e^{\langle \hat{c}_j, x \rangle - b_j} \right),$$

and finally define

$$c_j \stackrel{\text{def}}{=} \hat{c}_j - \nabla \hat{f}(0), \quad j = 1, \dots, m.$$

Note that by construction

$$\nabla f(0) \stackrel{(5.2)}{=} \frac{1}{\sum_{j=1}^m e^{-b_j}} \sum_{j=1}^m e^{-b_j} (\hat{c}_j - \nabla \hat{f}(0)) = 0,$$

so the unique minimizer of our test function (5.1) is $x^* = 0$. The starting point x_0 for all methods is the same and generated randomly from the uniform distribution on the standard Euclidean sphere of radius $1/n$ (this choice is motivated by (4.28)).

Thus, our test function (5.1) has three parameters: the dimension n , the number m of linear functions, and the regularization coefficient γ . Let us present computational results for different values of these parameters. The termination criterion for all methods is $f(x_k) - f(x^*) \leq \epsilon(f(x_0) - f(x^*))$.

In the tables below, for each method, we display the number of iterations until its termination. The minus sign (–) means that the method has not been able to achieve the required accuracy after $1000n$ iterations.

Table 1: $n = m = 50, \gamma = 1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	79	4	4	3	45	35	34
10^{-3}	1812	777	57	18	342	57	52
10^{-5}	5263	1866	107	29	738	72	58
10^{-7}	8873	2836	158	39	917	83	63
10^{-9}	12532	3911	203	48	1028	93	67

Table 2: $n = m = 50, \gamma = 0.1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	76	4	4	3	44	33	33
10^{-3}	2732	1278	78	23	512	70	56
10^{-5}	29785	12923	254	57	3850	126	72
10^{-7}	–	23245	346	74	6794	169	81
10^{-9}	–	32441	381	79	8216	204	87

Table 3: $n = m = 250, \gamma = 1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	444	4	4	3	214	158	157
10^{-3}	10351	4743	98	21	3321	264	251
10^{-5}	73685	31468	288	55	15637	350	274
10^{-7}	159391	58138	450	82	21953	413	296
10^{-9}	249492	85218	627	110	25500	464	314

Table 4: $n = m = 250, \gamma = 0.1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	442	4	4	3	209	155	155
10^{-3}	9312	4175	91	21	2686	258	251
10^{-5}	207978	102972	488	87	60461	556	346
10^{-7}	—	—	1003	170	147076	792	391
10^{-9}	—	—	1407	233	212100	976	419

We see that all quasi-Newton methods outperform the gradient method and demonstrate superlinear convergence (from some moment, the difference in the number of iterations between successive rows in the table becomes smaller and smaller). Among quasi-Newton methods (both the standard and the greedy ones), SR1 is always better than BFGS, while DFP is significantly worse than the other two. At the first few iterations, the greedy methods loose to the standard ones, but later they catch up. However, the classical SR1 method always remains the best. Nevertheless, the greedy methods are quite competitive.

Now let us look at the quality of Hessian approximations, produced by the quasi-Newton methods. In the tables below, we display the desired accuracy ϵ vs the final Hessian approximation error (defined as the operator norm of $G_k - \nabla^2 f(x_k)$, measured with respect to $\nabla^2 f(x_k)$). We look at the same problems as in Table 1 and Table 3.

Table 5: $n = m = 50, \gamma = 1$

ϵ	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-0}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$
10^{-1}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$2.7 \cdot 10^3$	$1.5 \cdot 10^3$	$1.5 \cdot 10^3$
10^{-3}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.2 \cdot 10^3$	$1.2 \cdot 10^1$	$3.8 \cdot 10^0$
10^{-5}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$2.1 \cdot 10^2$	$7.2 \cdot 10^0$	$2.6 \cdot 10^0$
10^{-7}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$9.1 \cdot 10^1$	$5.6 \cdot 10^0$	$2.2 \cdot 10^0$
10^{-9}	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$1.6 \cdot 10^3$	$5.2 \cdot 10^1$	$4.1 \cdot 10^0$	$1.8 \cdot 10^0$

Table 6: $n = m = 250, \gamma = 1$

ϵ	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-0}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$
10^{-1}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$7.1 \cdot 10^4$	$3.8 \cdot 10^4$	$3.9 \cdot 10^4$
10^{-3}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$6.8 \cdot 10^4$	$6.6 \cdot 10^1$	$1.7 \cdot 10^1$
10^{-5}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$9.4 \cdot 10^3$	$3.7 \cdot 10^1$	$1.2 \cdot 10^1$
10^{-7}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$3.1 \cdot 10^3$	$2.8 \cdot 10^1$	$9.7 \cdot 10^0$
10^{-9}	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$4.1 \cdot 10^4$	$1.7 \cdot 10^3$	$2.2 \cdot 10^1$	$7.3 \cdot 10^0$

As we can see from these tables, for standard quasi-Newton methods the Hessian approximation error always stays at the initial level. In contrast, for the greedy ones, it decreases relatively fast (especially for GrBFGS and GrSR1). Note also that sometimes the initial residual slightly increases at the first several iterations (which is noticeable only for GrDFP). This happens due to the fact that the objective function is non-quadratic, and we apply the correction strategy.

Note that in all the above tests we have used the same values for the parameters n and m . Let us briefly illustrate what happens when, for example, $m > n$.

Table 7: $n = 50, m = 100, \gamma = 0.1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	84	4	4	3	46	37	37
10^{-3}	897	316	32	11	183	53	52
10^{-5}	2421	833	67	19	334	63	58
10^{-7}	4087	1304	98	25	423	71	62
10^{-9}	5810	1859	132	32	473	78	66

Table 8: $n = 50, m = 200, \gamma = 0.1$

ϵ	GM	DFP	BFGS	SR1	GrDFP	GrBFGS	GrSR1
10^{-1}	108	4	4	3	45	46	46
10^{-3}	479	101	17	7	97	53	52
10^{-5}	1059	338	39	12	154	62	59
10^{-7}	1817	615	62	18	206	67	64
10^{-9}	2659	807	81	21	234	73	68

Comparing these tables with Table 2, we see that, with the increase of m , all the methods generally terminate faster. However, the overall picture is still the same as before. The results for $m < n$ are similar, so we do not include them.

Finally, let us present the results for the *randomized* version of scheme (4.18), in which, at every step, we select the update direction uniformly at random from the standard Euclidean sphere:

$$u_k \sim \text{Unif}(\mathcal{S}^{n-1}), \quad (5.5)$$

where $\mathcal{S}^{n-1} \stackrel{\text{def}}{=} \{x \in \mathbb{R}^n : \|x\| = 1\}$. We call the corresponding methods RaDFP, RaBFGS and RaSR1.

Table 9: $n = m = 50, \gamma = 1$

ϵ	RaDFP	RaBFGS	RaSR1
10^{-1}	35	29	34
10^{-3}	566	102	64
10^{-5}	1156	125	77
10^{-7}	1481	142	85
10^{-9}	1698	156	91

Table 10: $n = m = 250, \gamma = 1$

ϵ	RaDFP	RaBFGS	RaSR1
10^{-1}	261	144	158
10^{-3}	4276	366	287
10^{-5}	19594	517	346
10^{-7}	33293	619	376
10^{-9}	41177	698	396

It is instructive to compare these tables with Table 1 and Table 3, which contain the results for the greedy methods on the same problems. We see that the randomized methods are slightly slower than the greedy ones. However, the difference is not really significant, and, what is especially interesting, the randomized methods do not lose superlinear convergence.

6 Discussion

We have presented greedy quasi-Newton methods, that are based on the updating formulas from the Broyden family and use greedily selected basis vectors for updating Hessian approximations. For these methods, we have established explicit non-asymptotic rate of local superlinear convergence for the iterates and also a linear convergence for the deviations of Hessian approximations from the correct Hessians.

Clearly, there is a number of open questions. First, at every iteration, our methods need to compute the greedily selected basis vector. This requires additional information beyond just the gradient of the objective function (such as the diagonal of the Hessian). However, many problems, that arise in applications, possess certain structure (separable, sparse, etc.), for which the corresponding computations have a cost similar to that of the gradient evaluation (such as the test function in our experiments). Nevertheless, ideally it is desirable to get rid of the necessity in this auxiliary information at all. A natural idea might be to replace the greedy strategy with a *randomized* one. Indeed, as can be seen from our experiments, the corresponding scheme (4.18), (5.5) demonstrate almost the same performance as the greedy one. Therefore, one can expect that it should be possible to establish similar theoretical results about its superlinear convergence. Nevertheless, at the moment, we do not know how to do this. Although it is not difficult to show that, in terms of *expectations*, the randomized strategy still preserves the linear convergence of Hessian approximations (see [25]), it is not clear how to proceed after this in proving the superlinear convergence of the iterates, even in the quadratic case. The main difficulty, arising in the analysis, is that, at some moment, one needs to take the expectation of the product of random variables with known expectations, but the random variables themselves are *non-independent*.

Second, we have analyzed together a whole class of Hessian approximation updates by essentially upper bounding all its members via the worst one—DFP. Thus, all the efficiency guarantees, that we have established, might be too pessimistic for other members of this class such as BFGS, and especially SR1. Indeed, in our experiments, we have seen that the convergence properties of these three methods might differ quite significantly. It is

therefore desirable to refine our current analysis and obtain separate estimates for different updates.

Third, note that our current results do not prove anything about the rate of superlinear convergence of the standard quasi-Newton methods. Of course, it would be interesting to obtain the corresponding estimates and compare them to the ones, that we have established in this work.

Finally, apart from the quadratic case, we have not addressed at all the question of global convergence.

In any case, we believe that the ideas and the theoretical analysis, presented in this paper, will be useful for future advances in the theory of quasi-Newton methods.

References

- [1] W. Davidon. Variable metric method for minimization. Argonne National Laboratory Research and Development Report 5990 (1959).
- [2] C. Broyden. Quasi-Newton methods and their application to function minimization. *Mathematics of Computation*, **21**(99), 368-381 (1967).
- [3] R. Fletcher and M. Powell. A rapidly convergent descent method for minimization. *Computer Journal*, **6**(2), 163-168 (1963).
- [4] C. Broyden. The convergence of a class of double-rank minimization algorithms: 1. General considerations. *IMA Journal of Applied Mathematics*, **6**(1), 76-90 (1970).
- [5] C. Broyden. The convergence of a class of double-rank minimization algorithms: 2. The new algorithm. *IMA Journal of Applied Mathematics*, **6**(3), 222-231 (1970).
- [6] R. Fletcher. A new approach to variable metric algorithms. *Computer Journal*, **13**(3), 317-322 (1970).
- [7] D. Goldfarb. A family of variable-metric methods derived by variational means. *Mathematics of Computation*, **24**(109), 23-26 (1970).
- [8] D. Shanno. Conditioning of quasi-Newton methods for function minimization. *Mathematics of Computation*, **24**(111), 647-656 (1970).
- [9] M. Powell. On the convergence of the variable metric algorithm. *IMA Journal of Applied Mathematics*, **7**(1), 21-36 (1971).
- [10] C. Broyden, J. Dennis, and J. Moré. On the local and superlinear convergence of quasi-Newton methods. *IMA Journal of Applied Mathematics*, **12**(3), 223-245 (1973).
- [11] J. Dennis and J. Moré. A characterization of superlinear convergence and its application to quasi-Newton methods. *Mathematics of Computation*, **28**(126), 549-560 (1974).
- [12] J. Dennis and J. Moré. Quasi-Newton methods, motivation and theory. *SIAM Review*, **19**(1), 46-89 (1977).
- [13] B. Pshenichnyi and I. Danilin. Numerical methods in extremal problems. *Mir Publishers* (1978).

- [14] A. Stachurski. Superlinear convergence of Broyden’s bounded θ -class of methods. *Mathematical Programming*, **20**(1), 196-212 (1981).
- [15] A. Griewank and P. Toint. Local convergence analysis for partitioned quasi-Newton updates. *Numerische Mathematik*, **39**(3), 429-448 (1982).
- [16] R. Byrd, J. Nocedal, and Y. Yuan. Global convergence of a class of quasi-Newton methods on convex problems. *SIAM Journal on Numerical Analysis*, **24**(5), 1171-1190 (1987).
- [17] R. Byrd and J. Nocedal. A tool for the analysis of quasi-Newton methods with application to unconstrained minimization. *SIAM Journal on Numerical Analysis*, **26**(3), 727-739 (1989).
- [18] J. Engels and H. Martínez. Local and superlinear convergence for partially known quasi-Newton methods. *SIAM Journal on Optimization*, **1**(1), 42-56 (1991).
- [19] Y. Nesterov and A. Nemirovskii. Interior-point polynomial algorithms in convex programming. *SIAM*, **13** (1994).
- [20] H. Yabe and N. Yamaki. Local and superlinear convergence of structured quasi-Newton methods for nonlinear optimization. *Journal of the Operations Research Society of Japan*, **39**(4), 541-557 (1996).
- [21] Z. Wei, G. Yu, G. Yuan, and Z. Lian. The superlinear convergence of a modified BFGS-type method for unconstrained optimization. *Computational Optimization and Applications*, **29**(3), 315-332 (2004).
- [22] J. Nocedal and S. Wright. Numerical optimization. *Springer Science & Business Media* (2006).
- [23] H. Yabe, H. Ogasawara, and M. Yoshino. Local and superlinear convergence of quasi-Newton methods based on modified secant conditions. *Journal of Computational and Applied Mathematics*, **205**(1), 617-632 (2007).
- [24] A. Lewis and M. Overton. Nonsmooth optimization via quasi-Newton methods. *Mathematical Programming*, **141**(1-2), 135-163 (2013).
- [25] R. Gower and P. Richtárik. Randomized quasi-Newton updates are linearly convergent matrix inversion algorithms. *SIAM Journal on Matrix Analysis and Applications*, **38**(4), 1380-1409 (2017).
- [26] A. Mokhtari, M. Eisen, and A. Ribeiro. IQN: An incremental quasi-Newton method with local superlinear convergence rate. *SIAM Journal on Optimization*, **28**(2), 1670-1698 (2018).
- [27] Y. Nesterov. Lectures on convex optimization. *Springer*, **137** (2018).
- [28] N. Doikov and Y. Nesterov. Minimizing uniformly convex functions by cubic regularization of Newton method. *arXiv*, 1905.02671 (2019)
- [29] T. Sun and Q. Tran-Dinh. Generalized self-concordant functions: a recipe for Newton-type methods. *Mathematical Programming*, **178**, 145-213 (2019).
- [30] W. Gao and D. Goldfarb. Quasi-Newton methods: superlinear convergence without line searches for self-concordant functions. *Optimization Methods and Software*, **34**(1), 194-217 (2019).