
Improving Policy-Constrained Kidney Exchange via Pre-Screening

Anonymous Author(s)

Affiliation

Address

email

Abstract

In barter exchanges, participants swap goods with one another without exchanging money; these exchanges are often facilitated by a central clearinghouse, with the goal of maximizing the aggregate quality (or number) of swaps. Barter exchanges are subject to many forms of uncertainty—in participant preferences, the feasibility and quality of various swaps, and so on. Our work is motivated by kidney exchange, a real-world barter market in which patients in need of a kidney transplant swap their willing living donors, in order to find a better match. Modern exchanges include 2- and 3-way swaps, making the kidney exchange *clearing problem* NP-hard. Planned transplants often *fail* for a variety of reasons—if the donor organ is rejected by the recipient’s medical team, or if the donor and recipient are found to be medically incompatible. Due to 2- and 3-way swaps, failed transplants can “cascade” through an exchange; one US-based exchange estimated that about 85% of planned transplants failed in 2019. Many optimization-based approaches have been designed to avoid these failures; however most exchanges cannot implement these methods, due to legal and policy constraints. Instead, we consider a setting where exchanges can *query* the preferences of certain donors and recipients—asking whether they would accept a particular transplant. We characterize this as a two-stage decision problem, in which the exchange program (a) queries a small number of transplants before committing to a matching, and (b) constructs a matching according to fixed policy. We show that selecting these edges is a challenging combinatorial problem, which is non-monotonic and non-submodular, in addition to being NP-hard. We propose both a greedy heuristic and a Monte Carlo tree search, which outperforms previous approaches, using experiments on both synthetic data and real kidney exchange data from the United Network for Organ Sharing.

1 Introduction

We consider a multi-stage decision problem in which a decision-maker uses a fixed *policy* to solve a hard (stochastic) problem. Before using the policy, the decision-maker can first *measure* some of the uncertain problem parameters—in a sense, guiding the policy toward a better solution. Our primary motivation is kidney exchange, a process where patients in need of a kidney transplant swap their (willing) living donors, in order to find a better match. Many government-run kidney exchanges match patients and donors using a *matching algorithm* that follows strict policy guidelines [8]; this matching algorithm is often written into law or policy, and is not easily modified. Modern kidney exchanges use both cyclical swaps and chain-like structures (initiated by an unpaired altruistic donor) [23], and identifying the max-size or max-weight set of transplants is both NP- and APX-hard [1, 7].

In kidney exchange—as in many resource allocation settings—information used by the decision-maker is subject to various forms of uncertainty. Here we are primarily concerned with uncertainty in the *feasibility* of potential transplants: if a donor is matched with a potential recipient, will the transplant actually occur? Planned transplants may *fail* for a variety of reasons: for example, medical

39 testing may reveal that the donor and recipient are incompatible (a *positive crossmatch*); the recipient
40 or their medical team may reject a donor organ in order to wait for a better match; or the donor
41 may decide to donate elsewhere before the exchange is planned. Failed transplants are especially
42 troublesome in kidney exchange, due to the cycle and chain structures used: for example, suppose
43 that a cyclical swap is planned between three patient/donor pairs; if any one of the planned transplants
44 fails, then none of the other transplants in that cycle can occur. Unfortunately, it is quite common for
45 planned transplants to fail. For example, the United Network for Organ Sharing (UNOS¹) estimates
46 that in FY2019, about 85% of their planned kidney transplants failed [18].

47 Various matching algorithms have been proposed that aim to mitigate transplant failures (for exam-
48 ple, using stochastic optimization [15, 3], robust optimization [21], or conditional value at risk [6]).
49 However, implementing these strategies would require modifying fielded matching algorithms—which
50 in many cases would require changing law or policy. One way to avoid failures without modifying
51 the matching algorithm is to *pre-screen* potential transplants [18, 9, 10], by communicating with the
52 recipients’ medical team and possibly using additional medical tests. Pre-screening transplants is
53 costly, as it requires scarce time and resources. Furthermore, there are often many thousand potential
54 transplants in any given exchange; selecting which transplants to screen is not easy.

55 In this paper we investigate methods for selecting a limited number of transplants to pre-screen,
56 in order to “guide” the matching algorithm to a better outcome. We formalize this as a multistage
57 stochastic optimization problem, and we consider both an *offline* setting (where screenings are
58 selected all at once), and an *online* setting (where screenings are selected sequentially).

59 **Related Work.** While kidney exchange is known to be a hard packing problem, several algorithms
60 exist that are scalable in practice, and are used by fielded exchanges [14, 3, 19]. Prior work has
61 addressed potential transplant failures; our model is inspired by Dickerson et al. [15]. Pre-screening
62 potential transplants has also been addressed in prior work ([10, 22], and § 5.1 of [12]), and our model
63 is similar to stochastic matching and stochastic k -set packing [5]. However there are substantial
64 differences between these models and ours: (a) many prior approaches assume that a large number
65 of transplants may be pre-screened [10, 22]—on the order of one for each patient in the exchange;
66 we assume far fewer screenings are possible; (b) prior work often assumes a *query-commit* setting—
67 where successfully pre-screened transplants *must* be matched. Instead we assume that non-screened
68 transplants may also be matched—which more-accurately represents the way that modern exchanges
69 operate; (c) most prior work assumes that transplants that pass pre-screening are guaranteed to result
70 in a transplant. In reality, transplants often fail after pre-screening, a fact reflected in our model.

71 One of our approaches is based on *Monte Carlo Tree Search* (MCTS), which allows efficient explo-
72 ration of intractably large decision trees. While MCTS is primarily associated with Markov decision
73 processes and game-playing [11], it has been used successfully for combinatorial optimization [16].
74 We use a version of MCTS, Upper Confidence Bounds for Trees (UCT), which balances exploration
75 and exploitation by treating each tree node as a multi-armed bandit problem [4, 17].

76 Our Contributions

- 77 1. (§ 2) We formalize the *policy-constrained edge query problem*: where a decision-maker (such
78 as a kidney exchange program) selects a set of potential edges (potential transplants) to pre-
79 screen, prior to constructing a final packing (a set of transplants) using a fixed algorithm. This
80 model generalizes existing models in the literature, as edge failure probabilities depend on
81 whether or not the edge is pre-screened. Further, we allow for context-specific constraints,
82 such as those imposed by public policy or the particular hospital or exchange.
- 83 2. (§ 3) We prove that when the decision-maker uses a max-weight packing policy (the
84 most common choice among fielded exchanges), the edge query problem is both non-
85 monotonic and non-submodular in the set of queried edges. Despite these worst-case
86 findings we show that this problem is nearly monotonic for real and synthetic data, and
87 simple algorithms perform quite well. On the other hand, when the decision-maker uses
88 a *failure-aware* (stochastic) packing policy, the edge query problem becomes monotonic
89 under mild assumptions.
- 90 3. (§ 4) We conduct numerical experiments on both simulated and real exchange data from the
91 United Network for Organ Sharing (UNOS). We demonstrate that our methods substantially
92 outperform prior approaches and a randomized baseline.

¹UNOS is the organization tasked with overseeing organ transplantation in the US: <https://unos.org/>.

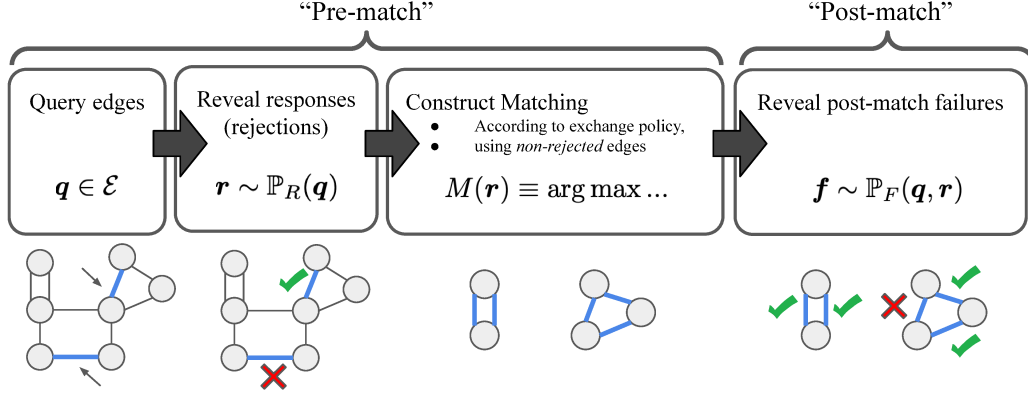


Figure 1: Single-stage edge selection: First, edges are selected to be queried, and responses revealed. Then, a final matching is constructed according to the exchange’s matching policy. Finally, the post-match edge failures are revealed.

2 The Policy-Constrained Edge Query Problem

Kidney exchanges are represented by a graph $G = (E, V)$ where vertices V represent (incompatible) patient-donor pairs, and non-directed donors (NDDs) who are willing to donate without receiving a kidney in return. Directed edges $e \in E$ between vertices represent potential transplants from the donor of one vertex to the patient of another. Edge weights represent the “utility” of an edge, and are typically set by exchange policy. Solutions to a kidney exchange problem (henceforth, *matchings*) consist of both directed *cycles* on G containing only patient-donor pairs, and directed *chains* beginning with an NDD and passing through one or more pairs; see Appendix A for an example exchange graph. Each vertex may participate in only one edge in a matching—as each vertex can donate and receive at most one kidney.

Vectors are denoted in bold, and are indexed by either cycles or edges: $\mathbf{y}_e$ indicates the element of $\mathbf{y}$ corresponding to edge e , and $\mathbf{x}_c$ is the element of $\mathbf{x}$ corresponding to cycle c . Our notation uses a *cycle-chain* representation for matchings²: let $\mathcal{C}$ represent cycles and chains in G , where each cycle and chain corresponds to a list of edges; as is standard in modern exchanges, we assume that cycles and chains are limited in length. Matchings are expressed as a binary vector $\mathbf{x} \in \{0, 1\}^{|\mathcal{C}|}$, where $\mathbf{x}_c = 1$ if cycle/chain c is in the matching, and 0 otherwise. Let w_c be the weight of cycle/chain c (the sum of c ’s edge weights). Let $\mathcal{M}$ denote the set of *feasible matching*—that is, the set of vertex-disjoint cycles and chains on G . The total weight of a matching is simply the summed weights of all its constituent cycles and chains: $\sum_{c \in \mathcal{C}} \mathbf{x}_c w_c$. We denote *sets* of edges using binary vectors, where $\mathbf{q} \in \{0, 1\}^{|E|}$ represents the set of all edges with $q_e = 1$.

In the remainder of this paper we refer to pre-screening a transplant as *querying an edge*, in order to be consistent with the literature.

Selecting Edge Queries. Our setting consists of two phases (see Figure 1): during *pre-match*, the decision-maker selects edges to query, and each queried edge is either accepted or rejected; then the decision-maker constructs a matching using a fixed policy. During *post-match*, each match edge either fails (no transplant) or succeeds (the transplant proceeds). We consider two versions of the pre-match phase: in the *single-stage* version, the decision-maker selects all queries before observing edge responses (accept/reject); in the *multi-stage* version, one edge is selected at a time and responses are observed immediately.

Unlike most prior work, edges in our model may fail during both the pre- and post-match phase. For example, suppose the decision-maker queries an edge from a 60-year-old non-directed donor, to a 35-year-old recipient; if the recipient or their medical team rejects the elderly donor and decides to wait for a younger donor, this is a pre-match rejection. Instead suppose the edge is not queried, and it is included in the final matching; if medical screening reveals that the patient and donor are incompatible, this is a post-match failure. We refer to pre-match failures as *rejections* and post-match failures as *failures*; however we make no assumption about their cause. We represent potential failures

²Our experiments use the position-indexed formulation, which is more compact and equivalent [14].

and rejections using binary random variables: $\mathbf{r} \in \{0, 1\}^{|E|}$ denotes pre-match rejections, where $\mathbf{r}_e = 1$ if e is queried and rejected, and 0 otherwise ($\mathbf{r}_e = 0$ for all non-queried edges). Similarly $\mathbf{f} \in \{0, 1\}^{|E|}$ denotes post-match failures, where $\mathbf{f}_e = 1$ if edge e fails post-match, and 0 otherwise. We assume that the distribution of rejections $\mathbf{r} \sim \mathbb{P}_R(\mathbf{q})$ is known, and depends on $\mathbf{q}$; we assume the distribution of failures $\mathbf{f} \sim \mathbb{P}_F(\mathbf{q}, \mathbf{r})$ is known, and depends on both $\mathbf{q}$ and $\mathbf{r}$.

Rejections and failures impact the matching through the *weight* of each cycle and chain. If any cycle edge fails, then *no* transplants in the cycle can proceed; if a chain edge fails, then all edges *following* it cannot proceed.³ Suppose we observe failures $\mathbf{f}$; the *final matching weight* of c is

$$F(c, \mathbf{y}) \equiv \begin{cases} \sum_{e \in c} w_e & \text{if } \sum_{e \in c} \mathbf{y}_e = 0 \\ 0 & \text{if } c \text{ is a cycle and } \sum_{e \in c} \mathbf{y}_e > 0 \\ \sum_{e \in c'} w_e & \text{if } c \text{ is a chain, where } c' \text{ includes all edges up to the first failed edge.} \end{cases}$$

Thus the *post-match expected weight* of matching $\mathbf{x}$, due to both rejections $\mathbf{r}$ and failures $\mathbf{f}$, is

$$W(\mathbf{x}; \mathbf{q}, \mathbf{r}) \equiv \mathbb{E}_{\mathbf{f} \sim \mathbb{P}_F(\mathbf{q}, \mathbf{r})} \left[\sum_{c \in \mathcal{C}} \mathbf{x}_c F(c, \mathbf{r} + \mathbf{f}) \right].$$

Matching Policy In this paper we assume that the final matching is constructed using a fixed matching policy, which uses only *non-rejected* edges; we denote this policy by $M(\mathbf{r})$. We focus primarily on the *max-weight* policy $M^{\text{MAX}}(\cdot)$, which is used by most fielded exchanges, and the *failure-aware* policy $M^{\text{FA}}(\cdot)$, which maximizes the expected post-match weight [15]:

$$M^{\text{MAX}}(\mathbf{r}) \in \arg \max_{\mathbf{x} \in \mathcal{M}} \sum_{c \in \mathcal{C}} \mathbf{x}_c F(c, \mathbf{r}), \quad M^{\text{FA}}(\mathbf{r}) \in \arg \max_{\mathbf{x} \in \mathcal{M}(\mathbf{r})} \mathbb{E}_{\mathbf{f} \sim \mathbb{P}_F(\mathbf{q}, \mathbf{r})} \left[\sum_{c \in \mathcal{C}} \mathbf{x}_c F(c, \mathbf{r} + \mathbf{f}) \right].$$

Next we formalize the *edge selection problem*—the main focus of this paper. We denote by $\mathcal{E}$ the set of “legal” edge subsets, subject to exchange-specific constraints; we assume that $\mathcal{E}$ is a matroid with ground set E . For example, the decision-maker may limit the number of queries issued to any one medical team (vertex in G) or transplant center (group of vertices). We aim to select an edge set $\mathbf{q} \in \mathcal{E}$ which maximizes the *expected weight* of the final matching. These edges are selected using only the distribution of future rejections and failures; we take a *stochastic optimization* approach, maximizing the expected outcome over this uncertainty.

Single-Stage Setting. The single-stage policy-constrained edge selection problem (henceforth, the *edge selection problem*) is expressed as

$$\max_{\mathbf{q} \in \mathcal{E}} V^S(\mathbf{q}), \quad \text{with } V^S(\mathbf{q}) \equiv \mathbb{E}_{\mathbf{r} \sim \mathbb{P}_R(\mathbf{q})} [W(M(\mathbf{r}); \mathbf{q}, \mathbf{r})], \quad (1)$$

where, $M(\mathbf{r})$ denotes the matching policy after observing rejections $\mathbf{r}$, and $W(\mathbf{x}; \mathbf{q}, \mathbf{r})$ denotes the post-match expected weight of matching $\mathbf{x}$. Exact evaluation of $V^S(\mathbf{q})$ is often intractable, as the support of $\mathbb{P}_R(\mathbf{q})$ grows exponentially in $|\mathbf{q}|$. In experiments we approximate $V^S(\mathbf{q})$ using sampling, and these approximations converge for a moderate number of samples (see Appendix B).

Multistage Setting. In the multi-stage setting, edge rejections are observed immediately after each edge is queried. The multi-stage problem is expressed as

$$\max_{\mathbf{q}^1 \in \mathcal{E}_1} \mathbb{E}_{\mathbf{r}^1 \sim \mathbb{P}_R(\mathbf{q}^1)} \left[\max_{\mathbf{q}^2 \in \mathcal{E}_1} \mathbb{E}_{\mathbf{r}^2 \sim \mathbb{P}_R(\mathbf{q}^2)} \left[\dots \max_{\mathbf{q}^K \in \mathcal{E}_1} \mathbb{E}_{\mathbf{r}^K \sim \mathbb{P}_R(\mathbf{q}^K)} [W(M(\mathbf{r}); \mathbf{q}, \mathbf{r})] \dots \right], \quad (2)$$

where $\mathbf{q} \equiv \sum_{i=1}^K \mathbf{q}^i$ denotes all queried edges, $\mathbf{r} \equiv \sum_{i=1}^K \mathbf{r}^i$ denotes all rejections, and $\mathcal{E}_1 \subseteq \mathcal{E}$ be denotes the legal edge subsets containing only one edge. First, we observe that Problems 1 and 2 require evaluating a matching policy $M(\mathbf{r})$. In the case of kidney exchange, evaluating both the max-weight policy $M^{\text{MAX}}(\cdot)$ and the failure-aware policy $M^{\text{FA}}(\cdot)$ require solving NP-hard problems; thus Problems 1 and 2 are at least NP-hard as well.

However, regardless of matching policy, the question whether *edge selection* is hard. We observe that while these problems are difficult in principle, experiments (§ 4) show that they are easy in practice. Proofs of the following propositions can be found in Appendix D.

³This assumes that chains can be *partially* executed: for example, suppose that the 4th edge in a 10-edge chain fails; the first three edges can still be matched, and the post-failure chain weight sums only these three edges. Not all fielded exchanges use this policy: some exchanges cancel the entire chain if one of its edges fails.

Proposition 2.1. *With matching policy $M^{FA}(\cdot)$, the objective of Problem 1 is non-monotonic in the number of queried edges, even with independent edge distributions.*

In other words, querying additional edges can sometimes lead to a *worse* outcome. This is somewhat counter-intuitive; one might think that providing additional information to the matching policy would strictly improve the outcome. This is a worst-case result—and in fact our experiments demonstrate that querying edges almost always leads to a better final matching weight.

Proposition 2.2. *With matching policy $M^{MAX}(\cdot)$, the objective of Problem 1 is non-submodular in the set of queried edges.*

In other words, certain edges are *complementary* to each other—and querying complementary edges simultaneously can yield a greater improvement than querying them separately. Taken together, these propositions indicate that single-stage edge selection with matching policy $M^{MAX}(\cdot)$ is a challenging combinatorial optimization problem. On the other hand, using the failure-aware matching policy $M^{FA}(\cdot)$ allows us to avoid some of these issues.

Proposition 2.3. *With matching policy $M^{FA}(\cdot)$, and if all edges are independent, the objective of Problem 1 is monotonic in the set of queried edges.*

While Propositions 2.1 and 2.2 state that single-stage edge selection is challenging in the worst case, our computational results suggest that these problems are often easier on realistic exchanges.

3 Solving the Policy-Constrained Edge Query Problem

First we propose an exhaustive tree search which returns an optimal solution to Problem 1 given enough time. Building on this, we propose a Monte Carlo Tree Search algorithm and a simple greedy algorithm. Our multi-stage approaches are very similar to these, and can be found in Appendix E.

Our optimal exhaustive search uses a *search tree* where each tree node corresponds to an edge subset in $\mathcal{E}$. The children of node $\mathbf{q}$ correspond to any $\mathbf{q}' \in \mathcal{E}$ which are equivalent to the parent $\mathbf{q}$, but include one additional edge: $C(\mathbf{q}) \equiv \{(\mathbf{q} + \mathbf{q}') \mid \forall \mathbf{q}' \in \mathcal{E} : |\mathbf{q}'| = 1 \mid (\mathbf{q} + \mathbf{q}') \in \mathcal{E}\}$. We say that edge sets (or tree nodes) containing L edges are on the L^{th} level of the tree. We refer to nodes with no children as *leaf nodes*. Unlike other tree search settings, the optimal solution to Problem 1 may be at *any* node of the tree, not only leaf nodes; this is a consequence of non-monotonicity (see Proposition 2.1). The tree defined by root node $\mathbf{q} = \mathbf{0}$ and child function $C(\mathbf{q})$ contains all legal edge subsets in $\mathcal{E}$, when $\mathcal{E}$ is a matroid. Thus, *any* exhaustive tree search algorithm (such as depth-first search) will identify an optimal solution, given enough time and memory.

Of course exhaustive search is only tractable if $\mathcal{E}$ is small. Consider the class of *budgeted* edge sets $\mathcal{E}(\Gamma)$ used in our experiments: $\mathcal{E}(\Gamma) \equiv \{\mathbf{q} \in \{0, 1\}^{|E|} \mid |\mathbf{q}| \leq \Gamma\}$ (edge sets containing at most Γ edges). The number of edge sets in $\mathcal{E}(\Gamma)$ grows roughly exponentially in Γ and $|E|$, and is impossible to enumerate even for small graphs. Suppose a graph has 50 edges and we have an edge budget of five: there are over two million edge sets in $\mathcal{E}(5)$. Even small exchange graphs can have thousands of edges, and thus $\mathcal{E}(\Gamma)$ cannot be enumerated. Therefore, we propose search-based approach.

Monte Carlo Tree Search for Edge Selection (MCTS): We propose a tree-search algorithm for single-stage edge selection, MCTS, based on Monte Carlo Tree Search (MCTS), with the Upper Confidence for Trees (UCT) algorithm [17]. UCT aims to learn the value of tree nodes using repeated sampling and simulation. When the set of tree nodes is too large to enumerate UCT can use a huge amount of memory—by storing values for each visited node. To limit both memory use and runtime, we incrementally search the tree from a temporary root node. Beginning from the root (the the empty edge set), we use UCB sampling on the next L levels of nodes—where L is a small fixed integer. After a fixed time limit, sampling stops and we set the *new* root node to the current root’s best child according to its UCB estimate—using the method of [17]. This process repeats until we reach the final level of the search tree. Algorithm 1 gives a pseudocode description of MCTS, which uses Algorithm 2 as a submethod. While often successful, MCTS requires extensive training and parameter tuning. As a simpler alternative, we propose a greedy algorithm.

Single-Stage Greedy Algorithm: Greedy. Like MCTS, our greedy algorithm (Greedy) begins with the empty edge set as the root node, and iteratively searches deeper levels of the tree. However unlike MCTS, Greedy simply selects the child node with the greatest objective value in Problem 1—that is, *greedily* improving the objective value; see Appendix E for a pseudocode description.

ALGORITHM 1: MCTS: Tree Search for Single-Stage Edge Selection

(input) K : maximum size of any legal edge set
 (input) T : time limit per level
 (input) L : number of look-ahead levels

$q^R \leftarrow \mathbf{0}$ root node (no edges)
 $q^* \leftarrow \mathbf{0}$ the best visited node
 $V^* \leftarrow$ objective value of q^*
for $N = 1, \dots, K$ **do**
 $M \leftarrow \min\{N + L, K\}$
 $Q \leftarrow$ all nodes in levels N to M
 $U[q] \leftarrow 0 \forall q \in Q$ UCB value estimate
 $V[q] \leftarrow 0 \forall q \in Q$ objective value
 $N[q] \leftarrow 0 \forall q \in Q$ number of visits
 while *less than time T has passed* **do**
 $\text{Sample}(q^R, M)$
 $q^R \leftarrow \arg \max_{q \in C(q^R)} U[q]$
 Delete $U[\cdot]$, $V[\cdot]$, and $N[\cdot]$
return q^*

ALGORITHM 2: Sample: Sampling function used by MCTS

(input) q, M

$N[q] \leftarrow N[q] + 1$
 $V[q] \leftarrow$ objective of edge set q in Problem 1
if $V[q] > V^*$ **then**
 $q^* \leftarrow q, V^* \leftarrow V[q]$
if q has no children **then**
 return $V[q]$
if q has children **then**
 if $|q| < M$ **then**
 $q' \leftarrow \arg \max_{q \in C(q^R)} U[q] + \text{UCB}[q]$
 $U[q] \leftarrow U[q] + \text{Sample}(q', M)$
 else
 $q' \leftarrow$ a random descendent of q at any level
 $V' \leftarrow$ objective value of q' in Problem 1
 if $V' > V^*$ **then**
 $q^* \leftarrow q', V^* \leftarrow V'$
 $U[q] \leftarrow U[q] + V'$

4 Computational Experiments

We conduct a series of computational experiments using both synthetic data, and real kidney exchange data from UNOS; all code for these experiments is available online.⁴ In these experiments, “legal” edge sets are the budgeted edge sets defined as $\mathcal{E}(\Gamma) \equiv \{q \in \{0, 1\}^{|E|} \mid |q| \leq \Gamma\}$. In Sections 4.1 and 4.2 we present results in the single- and multi-stage edge selection settings, respectively. We use two types of data for these experiments:

Real Data. We use exchange graphs from the United Network for Organ Sharing (UNOS), representing UNOS match runs between 2010 and 2016. Some of these exchange graphs only have the trivial matching (no cycles or chains), or they have only one non-trivial matching. We ignore these graphs because the matching policy is a “constant” function (to return the one feasible matching) and edge queries cannot change the outcome. Removing these, we are left with 240 UNOS exchange graphs.

Synthetic Data. We generate random kidney exchange graphs based on directed Erdős-Rényi graphs defined using parameters N and p : let V be a fixed set of N vertices; for each pair of vertices (V_1, V_2) there is an edge from V_1 to V_2 with probability p , and an edge from V_2 to V_1 with probability p (independent of the edge from V_1 to V_2). Any vertices with no incoming edges are considered NDDs.

In these experiments edge rejections and failures are independently distributed for each edge e ; let P_R be the rejection probability, P_Q is the post-match success probability if e is queried/accepted, and P_N is the success probability if e is not queried. To simulate edge rejections and failures we use two synthetic edge distributions: *Simple* and *KPD*. In the *Simple* distribution, $P_R = 0.5$, $P_Q = 1$, and $P_N = 0.5$ for all edges. The *KPD* distribution is inspired by the fielded exchange setting from which we draw our real underlying compatibility graphs. According to UNOS, about 34% of all edges are rejected by a donor or recipient pre-match [18]; we draw P_R uniformly from $U(0.25, 0.43)$ for each edge. Edges ending in highly-sensitized patients (who are often less healthy and more likely to be incompatible) are considered high-risk; for these edges we draw P_Q from $U(0.2, 0.5)$ and P_N from $U(0.0, 0.2)$. For other edges we draw P_Q from $U(0.9, 1.0)$ and P_N from $U(0.8, 0.9)$.

4.1 Single-Stage Edge Selection Experiments

In this section we compare against the baseline of a max-weight matching *without* edge queries (using policy $M^{\text{MAX}}(\cdot)$). Let V_X be the objective⁵ of Problem 1 achieved by method X , we calculate Δ^{MAX} (the relative difference from baseline) as $\Delta^{\text{MAX}} \equiv (V_X - V^S(\mathbf{0}))/V^S(\mathbf{0})$. A value of $\Delta^{\text{MAX}} = 0$ means that method X did not improve over the baseline, a value of $\Delta^{\text{MAX}} = 1$ means that X achieved an objective 100% greater than the baseline, and so on.

⁴(link removed during review)

⁵All objective values are estimated using up to 1000 sampled rejection scenarios (see Appendix B), as it is intractable to evaluate the exact objective of large edge sets.

Table 1: Left: Optimality gap for Greedy, over 100 random graphs with $p = 0.01$ and various N , with edge budget $\Gamma = 3$; bottom row shows the maximum value of %OPT over all graphs. Right: Single-stage results on UNOS graphs using the variable IIAB edge budget (top rows), and the failure-aware method (bottom row). Columns P_X indicates the X^{th} percentile of $\Delta^{\max}$ over all UNOS graphs.

%OPT	Num. Graphs (out of 100)			Method	Simple edge dist.			KPD edge dist.		
	$N = 50$	$N = 75$	$N = 100$		P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}
[0, 0.1]	93	93	90	MCTS	0.40	0.67	1.11	0.05	0.46	3.56
(0, 1]	5	4	9	Greedy	0.42	0.72	1.13	0.01	0.49	3.56
(1, 2]	1	3	1	Random	0.00	0.10	0.45	-0.11	0.00	0.63
(2, 100]	1	0	0	IIAB	0.21	0.45	0.89	-0.24	0.14	2.66
Max %OPT	2.8	1.5	1.0	Fail-Aware	0.00	0.09	0.24	-0.25 [†]	0.00 [†]	2.17 [†]

First we investigate the *difficulty* of edge selection. Using random graphs, we compare Greedy to the *optimal* solution to Problem 1, found by exhaustive search (OPT). We generate three sets of 100 random graphs with $N = 50, 75$, and 100 vertices, and each with $p = 0.01$. For all graphs we run both OPT and Greedy with edge budget 3; we calculate the *optimality gap* of Greedy as $\%OPT \equiv 100 \times (V_{OPT} - V_{Greedy})/V_{OPT}$, where V_X denotes the objective achieved by method X . ($V_{OPT} > 0$ in all graphs used in these experiments.) If $\%OPT = 0$ then Greedy returns an optimal solution, and $\%OPT > 0$ means that Greedy is not optimal. Table 1 (left) shows the number of random graphs binned by %OPT, as well as the maximum %OPT over all graphs. For each N , Greedy returns an optimal solution for at least 90 of the 100 graphs; the *maximum* %OPT over all graphs is 2.8.

We also test Greedy on real UNOS graphs, using maximum budget 100. Figure 2a shows the median $\Delta^{\max}$ over all UNOS graphs, with shading between the 10th and 90th percentiles. Larger edge budgets almost never decrease the objective achieved by Greedy, and Greedy *never* produces a worse outcome than the baseline. Thus—in our setting—single-stage edge selection is effectively monotonic in our setting, and Greedy is an effective method.

Small Edge Budgets on Real UNOS Graphs We compare all methods on UNOS graphs, using smaller, more-realistic edge budgets from 1 to 10. For MCTS we use a 1-hour time limit per edge (Γ hours total). Figures 2b and 2d compare $\Delta^{\max}$ for MCTS, Greedy, and random edge selection, for the *Simple* and *KPD* edge distributions, respectively. We draw two conclusions from these results: (1) MCTS and Greedy produce almost identical results, further suggesting that Greedy is nearly optimal in our setting; (2) in our setting, edge selection is *effectively* monotonic, as $\Delta^{\max}$ almost never decreases. However Figure 2d gives an example of non-monotonicity for both Greedy and Random: in some cases, querying edges can lead to a *worse* outcome than querying no edges.

We also compare against two state-of-the-art approaches: the edge selection approach of [10] (IIAB), which uses a *variable* edge budget that depends on the graph structure; and the failure-aware matching policy of [15] (Fail-Aware), which does not query edges. We for the *KPD* distribution we use an approximation of Fail-Aware, which assumes a uniform edge failure probability. To our knowledge no scalable algorithm exists for the general Fail-Aware method. Table 1 (right) shows a comparison of all edge-selection methods—each using the variable edge budget of IIAB; the bottom row shows results for Fail-Aware.

4.2 Multi-Stage Edge Selection Experiments on UNOS Graphs

We run initial multi-stage edge selection experiments on a subset of 150 randomly chosen UNOS graphs. For each graph we test our multi-stage variants of MCTS and Greedy, and compare with a baseline of random edge selection; as before, MCTS uses a 1-hour training time per level. It is substantially harder to evaluate the multi-stage objective, as each edge edge-selection method changes depending on rejections observed in prior stages. Similarly, the MCTS search tree is orders of magnitude larger in the multi-stage setting: each node in tree corresponds to both an edge set *and* a rejection scenario (see Appendix E).

In these initial experiments we evaluate each method on 10 edge rejections *realizations* (only a small subset). We estimate $\Delta^{\max}$ for each method and each graph by averaging the final matching weight over all realizations. Figure 2c shows the results of these experiments.

[†]We use an approximation of Fail-Aware for the *KPD* dist.; *true* Fail-Aware should always have $\Delta^{\max} > 0$.

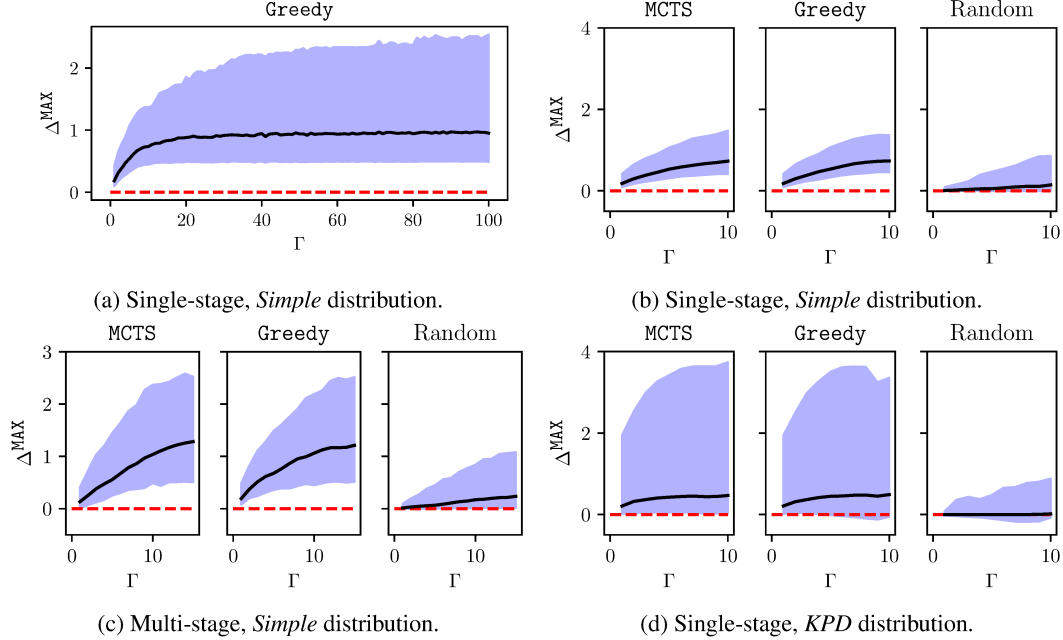


Figure 2: Results for UNOS graphs. Right: edge budget up to 10 for the *Simple* distribution (top) and the *KPD* distribution (bottom). Top-left: Greedy with edge budget up to 100, for the simple distribution. Bottom-left: multi-stage methods using the *Simple* distribution. In all plots, a solid line indicates median Δ^{MAX} over all UNOS graphs, and shading is between the 10th and 90th percentiles; a dotted line indicates the baseline.

These initial multi-stage results are quite similar to our single-stage results. However it is notable that the objective value in the multi-stage setting is somewhat higher than in the single-stage setting—even using the simple method Greedy. Further, this suggests that more can be gained by developing a more sophisticated multi-stage edge selection policy. We leave this for future work.

5 Conclusions and Future Research Directions

Many planned kidney exchange transplants *fail* for a variety of reasons; these failures greatly reduce the number of transplants that an exchange can facilitate, and increase the waiting time for many patients in need of a kidney. Avoiding transplant failures is a challenge, as exchanges are often constrained by policy and law in how they match patients and donors. We consider a setting where exchanges can *pre-screen* certain transplants, while still matching patients and donors using a fixed policy. We formalize a multi-stage optimization problem based on realistic assumptions about how transplants fail, and how exchanges match patients and donors; we emphasize that these important assumptions are not included in prior work. While this problem is challenging in theory, we show that it is much easier in practice—with computational experiments using both synthetic data and real data from the United Network for Organ Sharing. In experiments, we find that pre-screening even a small number of potential transplants (around 10) significantly increases the overall quality of the final match—by more than 100% of the original match weight.

Our initial study of the pre-screening problem suggests several areas for future work. First we assume that the distribution of transplant failures is known, when in reality only rough approximations of these distributions are available. Second, we assume that exchange participants (donors, recipients, hospitals) are not strategic. In reality, strategic behavior plays a substantial role in real exchanges [2]; we expect that participants might behave strategically when responding to pre-screening requests. Third, our model does not account for equitable treatment of different patients [20]. For example, it may be the case that pre-screening a transplant decreases the likelihood of the transplant being matched. That might disproportionately impact highly-sensitized patients, which are both sicker and more difficult to match than other patients.

Broader Impact

Patients with end-stage renal disease have only two options: receive a transplant, or undergo dialysis once every few days, for the rest of their lives. In many countries (including the US), these patients register for a deceased donor waiting list—and it can be months or years before they receive a transplant. Many of these patients have a friend or relative willing to donate a kidney, however many patients are incompatible with their corresponding donor. Kidney exchange allows patients to “swap” their incompatible donor, in order to find a *higher-quality* match, *more quickly* than a waiting list. Transplants allow patients a higher quality of life, and cost far less, than lifelong dialysis. About 10% of kidney transplants in the US are facilitated by an exchange.

Our aim in this paper is to understand the impact of transplant pre-screening on a mathematical formulation of kidney exchange. None of our experiments impact actual kidney exchange participants, however the methods we propose for pre-screening transplants could readily be applied by a real exchange, including nascent organ exchanges for other organs such as livers. Thus, we briefly discuss two potential impacts of our methods on a fielded exchange program:

- From a mathematical perspective, we show that pre-screening transplants can—in the worst case—negatively impact the final matching weight (or size) of an exchange (see Proposition 2.1). In other words, pre-screening can result in *fewer* patients receiving kidney transplants, and/or that the resulting transplants are of *lower quality*.
- In kidney exchange, some patients are harder to match than others. These hard-to-match (*highly sensitized*) patients are often sicker, and far less likely to find a compatible kidney donor, than other patients. As a consequence, highly sensitized patients are less likely to be matched by a kidney exchange algorithm [13, 20]; in fact, many exchanges (including the exchange run by UNOS) use policies to prioritize highly sensitized patients so they are not marginalized by a matching algorithm. There is a risk that, by algorithmically selecting transplants for pre-screening, we would further marginalize highly sensitized patients. This risk warrants further study, which we leave for future work.

References

- [1] D. Abraham, A. Blum, and T. Sandholm. Clearing algorithms for barter exchange markets: Enabling nationwide kidney exchanges. In *Proceedings of the ACM Conference on Electronic Commerce (EC)*, pages 295–304, 2007.
- [2] N. Agarwal, I. Ashlagi, E. Azevedo, C. R. Featherstone, and Ö. Karaduman. Market failure in kidney exchange. *American Economic Review*, 109(11):4026–70, 2019.
- [3] R. Anderson, I. Ashlagi, D. Gamarnik, and A. E. Roth. Finding long chains in kidney exchange using the traveling salesman problem. *Proceedings of the National Academy of Sciences*, 112(3):663–668, 2015.
- [4] P. Auer, N. Cesa-Bianchi, and P. Fischer. Finite-time analysis of the multiarmed bandit problem. *Machine learning*, 47(2-3):235–256, 2002.
- [5] N. Bansal, A. Gupta, J. Li, J. Mestre, V. Nagarajan, and A. Rudra. When LP is the cure for your matching woes: Improved bounds for stochastic matchings. *Algorithmica*, 63(4):733–762, 2012.
- [6] H. Bidkhori, J. P. Dickerson, K. Ren, and D. C. McElfresh. Kidney exchange with inhomogeneous edge existence uncertainty. In *Proceedings of the Conference on Uncertainty in Artificial Intelligence (UAI)*, forthcoming, 2020.
- [7] P. Biró, D. F. Manlove, and R. Rizzi. Maximum weight cycle packing in directed graphs, with application to kidney exchange programs. *Discrete Mathematics, Algorithms and Applications*, 1(04):499–517, 2009.
- [8] P. Biró, J. van de Klundert, D. Manlove, W. Pettersson, T. Andersson, L. Burnapp, P. Chromy, P. Delgado, P. Dworczak, B. Haase, et al. Modelling and optimisation in european kidney exchange programmes. *European Journal of Operational Research*, 2019.
- [9] A. Blum, A. Gupta, A. D. Procaccia, and A. Sharma. Harnessing the power of two crossmatches. In *Proceedings of the ACM Conference on Electronic Commerce (EC)*, pages 123–140, 2013.

- 363 [10] A. Blum, J. P. Dickerson, N. Haghtalab, A. D. Procaccia, T. Sandholm, and A. Sharma. Ignorance is almost bliss: Near-optimal stochastic matching with few queries. *Operations Research*, 2020. Earlier version appeared in the ACM Conference on Economics and Computation (EC), 2015.
- 364
- 365
- 366
- 367 [11] B. Bouzy. Associating shallow and selective global tree search with Monte Carlo for 9×9 Go. In *International Conference on Computers and Games*, pages 67–80. Springer, 2004.
- 368
- 369 [12] J. P. Dickerson. A unified approach to dynamic matching and barter exchange. Technical report, Doctoral Dissertation, Carnegie Mellon University, Pittsburgh, PA, 2016.
- 370
- 371 [13] J. P. Dickerson, A. D. Procaccia, and T. Sandholm. Price of fairness in kidney exchange. In *International Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, pages 1013–1020, 2014.
- 372
- 373
- 374 [14] J. P. Dickerson, D. Manlove, B. Plaut, T. Sandholm, and J. Trimble. Position-indexed formulations for kidney exchange. In *Proceedings of the ACM Conference on Economics and Computation (EC)*, 2016.
- 375
- 376
- 377 [15] J. P. Dickerson, A. D. Procaccia, and T. Sandholm. Failure-aware kidney exchange. *Management Science*, 65(4):1768–1791, 2019. Earlier version appeared in the ACM Conference on Economics and Computation (EC), 2013.
- 378
- 379
- 380 [16] B. Kartal, E. Nunes, J. Godoy, and M. Gini. Monte Carlo tree search with branch and bound for multi-robot task allocation. In *The IJCAI-16 workshop on autonomous mobile service robots*, volume 33, 2016.
- 381
- 382
- 383 [17] L. Kocsis and C. Szepesvári. Bandit based Monte-Carlo planning. In *European conference on machine learning*, pages 282–293. Springer, 2006.
- 384
- 385 [18] R. Leishman. Challenges in match offer acceptance in the optn kidney paired donation pilot program. In *INFORMS Annual Meeting*, 2019. Presentation in Session TB94 - Kidney Allocation & Exchange.
- 386
- 387
- 388 [19] D. Manlove and G. O’Malley. Paired and altruistic kidney donation in the UK: Algorithms and experimentation. *ACM Journal of Experimental Algorithmics*, 19(1), 2015.
- 389
- 390 [20] D. C. McElfresh and J. P. Dickerson. Balancing lexicographic fairness and a utilitarian objective with application to kidney exchange. *AAAI Conference on Artificial Intelligence (AAAI)*, 2018.
- 391
- 392 [21] D. C. McElfresh, H. Bidkhori, and J. P. Dickerson. Scalable robust kidney exchange. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pages 1077–1084, 2019.
- 393
- 394
- 395 [22] M. Molinaro and R. Ravi. Kidney exchanges and the query-commit problem. Manuscript, 2013.
- 396 [23] M. Rees, J. Kopke, R. Pelletier, D. Segev, M. Rutter, A. Fabrega, J. Rogers, O. Pankewycz, J. Hiller, A. Roth, T. Sandholm, U. Ünver, and R. Montgomery. A nonsimultaneous, extended, altruistic-donor chain. *New England Journal of Medicine*, 360(11):1096–1101, 2009.
- 397
- 398

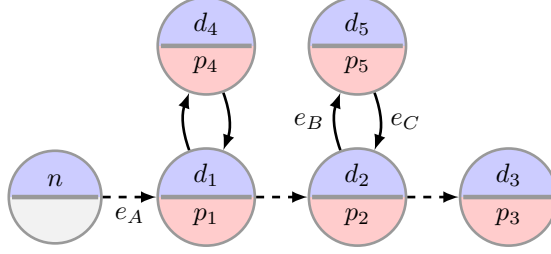


Figure 3: Sample exchange graph with a 3-chain (dashed edges) and two 2-cycles (solid edges). The NDD is denoted by n , and each patient (and associated donor) is denoted by p_i (d_i). If edge e_1 is not queried, or queried and *accepted*, then the chain may be included in the final matching. However if edge e_A is queried and *rejected*, then only the 2-cycles may be included in the final matching.

Edge budgets	$N = 10$	$N = 30$	$N = 50$	$N = 100$	$N = 1000$
1-10	0.0994	0.0568	0.0443	0.0316	0.00978
11-20	0.118	0.0689	0.0533	0.0373	0.0119
21-30	0.131	0.0758	0.0590	0.0418	0.0131
31-40	0.136	0.0787	0.0609	0.0436	0.0135
41-50	0.140	0.0812	0.0627	0.0442	0.014
51-60	0.146	0.0846	0.0653	0.0461	0.0147
61-70	0.152	0.0883	0.0680	0.0481	0.0153
71-80	0.162	0.093	0.0720	0.0514	0.0162
81-90	0.172	0.0995	0.0763	0.0544	0.0172
91-100	0.182	0.104	0.0815	0.0576	0.0180

Table 2: Median normalized standard deviation of the bootstrap mean, over 200 bootstrap samples for each sample size N , binned by edge budget.

A Kidney Exchange

B Estimating The Objective of Problem 1

DCM: @Michael can you review this for accuracy? I added and changed some words, which are hopefully correct

The objective of the single-stage edge selection problem requires evaluating all rejection scenarios $\mathbf{r} \sim \mathbb{P}_R(\mathbf{q})$, and the support of this distribution grows exponentially in the number of edges $|\mathbf{q}|$. In computational experiments, to estimate the objective of Problem 1, we sample up to 1000 scenarios from $\mathbb{P}_R(\mathbf{q})$. More explicitly: we *exactly* evaluate the objective of edge sets with fewer than 10 edges; for larger edge sets, we sample the objective using 1000 draws from $\mathbb{P}_R(\mathbf{q})$.

Using bootstrapping experiments we demonstrate that our sampling approach is sufficient to accurately estimate the true objective, even for large edge sets. For 152 UNOS graphs, we computed edge sets by running Greedy with edge budgets ranging from 1 to 100. For each edge set, we then sample a subset of $N \in \{10, 30, 50, 100, 1000\}$ rejection scenarios, with replacement, from the set of all sampled edge outcomes. For each edge set and choice of N we repeat 200 times and calculate the sample mean for each replication. We then compute the standard deviations of these bootstrap sample means to estimate the variance due to sampling. For each N , we calculate the mean sample standard deviation, normalized by the sample mean. Table 2 shows the median normalized standard deviation for all experiments under each N , with edge budgets aggregated into 10 bins. We find that with $N = 1000$ samples, the standard deviation was on average only about 2% of the overall mean value, even for large edge budgets.

Michael: do we need to cite bootstrapping? this is fairly ad-hoc.

Γ	MCTS			Greedy			Random			IIAB		
	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}
1	0.10	0.18	0.42	0.10	0.18	0.42	0.00	0.01	0.09			
2	0.17	0.30	0.65	0.17	0.29	0.64	0.00	0.02	0.16			
3	0.22	0.38	0.80	0.24	0.38	0.79	0.00	0.03	0.23			
4	0.27	0.45	0.91	0.28	0.47	0.91	0.00	0.05	0.39			
5	0.31	0.53	1.06	0.33	0.53	1.07	0.00	0.05	0.47			
6	0.35	0.58	1.16	0.38	0.60	1.18	0.00	0.07	0.54			
7	0.37	0.63	1.28	0.41	0.66	1.27	0.00	0.09	0.62			
8	0.39	0.67	1.34	0.43	0.70	1.33	0.00	0.11	0.74			
9	0.41	0.69	1.39	0.46	0.72	1.37	0.00	0.11	0.85			
10	0.41	0.73	1.48	0.46	0.73	1.37	0.01	0.15	0.86			
IIAB	0.40	0.67	1.11	0.42	0.72	1.13	0.00	0.10	0.45	0.21	0.45	0.89

Table 3: Single-stage results for the *simple* edge distribution on UNOS graphs: Δ^{MAX} for each method, and for each edge budget Γ . The 10th, 50th and 90th percentiles of Δ^{MAX} (P_{10} , P_{50} , and P_{90}) are reported across all UNOS graphs. Maximum value for each Γ is reported in bold.

Γ	MCTS			Greedy			Random			IIAB		
	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}	P_{10}	P_{50}	P_{90}
1	0.02	0.21	1.94	0.02	0.21	1.94	-0.03	0.00	0.08			
2	0.04	0.32	2.53	0.04	0.31	2.49	-0.04	0.00	0.36			
3	0.05	0.37	2.99	0.04	0.37	2.99	-0.06	0.00	0.43			
4	0.05	0.41	3.27	0.01	0.42	3.27	-0.07	0.00	0.38			
5	0.05	0.43	3.42	-0.01	0.45	3.50	-0.10	0.00	0.48			
6	0.05	0.44	3.57	-0.04	0.45	3.62	-0.14	0.00	0.66			
7	0.05	0.45	3.63	-0.07	0.47	3.63	-0.17	0.00	0.67			
8	0.03	0.44	3.64	-0.10	0.48	3.63	-0.18	0.00	0.78			
9	0.02	0.44	3.64	-0.12	0.45	3.25	-0.17	0.00	0.81			
10	0.04	0.47	3.74	-0.06	0.49	3.36	-0.07	0.02	0.88			
IIAB	0.05	0.46	3.56	0.01	0.49	3.56	-0.11	0.00	0.63	-0.24	0.14	2.66

Table 4: Single-stage results for the *KPD* edge distribution on UNOS graphs: Δ^{MAX} for each method, and for each edge budget Γ . The 10th, 50th and 90th percentiles of Δ^{MAX} (P_{10} , P_{50} , and P_{90}) are reported across all UNOS graphs. Maximum value for each Γ and each percentile is reported in bold.

C Additional Computational Experiment Results

D Proofs for Section 2

In the proofs of Proposition 2.1 and Proposition 2.2 we consider a setting where all edges’ pre-match rejections and post-match failures are i.i.d., where $P_R = 0.5$ is the pre-match rejection probability, $P_Q = 1.0$ is the post-match success probability if the edge is queried-and-accepted, and $P_N = 0.5$ is the success probability if e is not queried. That is, queried edges have rejection probability 0.5, accepted edges have zero failure probability, and non-queried edges have failure probability 0.5.

D.1 Proof of Proposition 2.1

(Proof by counterexample.) We provide an example where querying a single edge results in a *lower* objective value in Problem 1 (i.e., final expected matching weight) than querying no edges—when using the max-weight matching policy $M^{\text{MAX}}(\cdot)$.

Consider the exchange graph in Figure 4; edge (E, B) has weight 1.5, while all other edges have weight 1. First we consider the objective due to querying no edges, $V^S(\mathbf{0})$. In this case, no edges can be rejected pre-match, the max-weight matching includes cycle (C, D, F) (expected weight $3 \times (1/2)^3 = 3/8$) and cycle (A, B) (expected weight $2 \times (1/2)^2 = 1/2$), with total expected matching weight $7/8$. That is, $V^S(\mathbf{0}) = 7/8$.

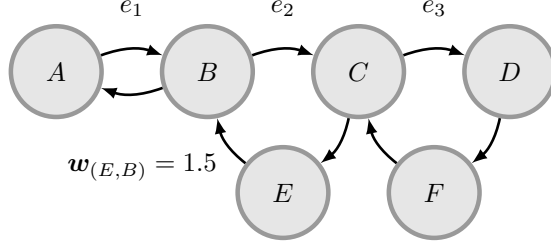


Figure 4: Exchange graph for Propositions 2.1 and 2.2. All edges have weight 1 except for edge (E, B) , which has weight 1.5.

Next consider the objective due to querying only edge $e_3 = (C, D)$, and let q' denote edge set $\{e_3\}$. With probability $1/2$, e_3 is rejected and cycle (B, C, E) is the max-weight matching – with expected weight $3.5/8$. With probability $1/2$, e_3 is accepted and the max-weight matching includes cycles (A, B) (with expected weight $1/2$) and (C, D, F) (with expected weight $3/4$); this matching has total expected weight $5/4$. Thus, $V^S(q) = 27/32 < 7/8 = V^S(0)$, which concludes the proof.

D.2 Proof of Proposition 2.2

(Proof by counterexample.) We provide an example where the objective value in Problem 1 (i.e., final expected matching weight) is non-submodular—when using the max-weight matching policy $M^{\max}(\cdot)$. We use the same rejection and failure distribution as in the proof of Proposition 2.1.

Consider the exchange graph in Figure 4; edge (E, B) has weight 1.5, while all other edges have weight 1. With some abuse of notation, we will denote by $V^S(\{e_a, \dots, e_N\})$ the objective of Problem 1 due to edge set $\{e_a, \dots, e_N\}$. Our counterexample for submodularity is that, for this graph,

$$V^S(X \cup \{e_1, e_2\}) + V^S(X) > V^S(X \cup \{e_1\}) + V^S(X \cup \{e_2\}),$$

with set $X \equiv \{e_3\}$. That is, the objective increase due to querying *both* edges e_1 and e_3 is greater than the combined increase due to querying both edges separately. Next we explicitly calculate each of the above terms.

$V^S(X) = V^S(\{e_3\})$. There are two cases to consider:

- e_3 is accepted, with probability $1/2$. The max-weight matching is cycles (A, B) and (C, D, F) , with expected weight $(1/2 + 3/4)$,
- e_3 is rejected, with probability $1/2$. The max-weight matching is cycle (B, C, E) , with expected weight $3.5/8$.

Thus, $V^S(X) = (1/2)(1/2 + 3/4) + (1/2)(3.5/8) = 27/32$.

$V^S(X \cup \{e_1\}) = V^S(\{e_1, e_3\})$. There are four cases to consider:

- e_1 and e_3 are accepted, with probability $1/4$. The max-weight matching is cycles (A, B) and (C, D, F) , with expected weight $(1 + 3/8)$,
- e_1 is rejected and e_3 is accepted, with probability $1/4$. The max-weight matching is cycle (B, C, E) , with expected weight $3.5/8$.
- e_1 is accepted and e_3 is rejected, with probability $1/4$. The max-weight matching is cycle (B, C, E) , with expected weight $3.5/8$.
- e_1 and e_3 are rejected, with probability $1/4$. The max-weight matching is cycle (B, C, E) , with expected weight $3.5/8$.

Thus the objective is $V^S(X \cup \{e_3\}) = (1/4)(1 + 3/8) + (3/4)(3.5/8) = 43/64$.

$V^S(X \cup \{e_2\}) = V^S(\{e_2, e_3\})$. There are three cases to consider

- e_3 is accepted, with probability $1/2$. The max-weight matching is cycles (A, B) and (C, D, F) , with expected weight $(1/2 + 3/4)$,
- e_3 is rejected and e_2 is accepted, with probability $1/4$. The max-weight matching is cycle (B, C, E) , with expected weight $3.5/4$,
- e_3 and e_2 are rejected, with probability $1/4$. The max-weight matching is cycle (A, B) , with expected weight $1/2$.

472 Thus the objective is $V^S(X \cup \{e_2\}) = (1/2)(1/2 + 3/4) + (1/4)(3.5/4) + (1/4)(1/2) = 31/32$.

473 $V^S(X \cup \{e_1, e_2\}) = V^S(\{e_1, e_2, e_3\})$. There are four cases to consider:

- 474 • e_1 and e_3 are accepted, with probability $1/4$. The max-weight matching is cycles (A, B)
475 and (C, D, F) , with expected weight $(1 + 3/4)$,
- 476 • e_1 is accepted and e_2 is rejected, with probability $1/4$ (the response from e_3 is irrelevant).
477 The max-weight matching is (A, B) and (C, D, F) , with expected weight $1 + 3/8$.
- 478 • e_1 is rejected and e_2 is accepted (the response from e_3 is irrelevant), with probability $1/4$.
479 The max-weight matching is cycle (B, C, E) , with expected weight $3.5/4$.
- 480 • e_1 and e_2 are rejected (the response from e_3 is irrelevant), with probability $1/4$. The
481 max-weight matching is cycle (C, D, F) , with expected weight $3/8$.

482 Thus the objective is $V^S(X \cup \{e_1, e_2\}) = (1/4)(1 + 3/4) + (1/4)(1 + 3/8) + (1/4)(3.5/4) +$
483 $(1/4)(3/8) = 35/32$.

484 Finally, we have:

$$V^S(X \cup \{e_1, e_2\}) + V^S(X) = 35/32 + 27/32 = 1.9375$$

485 and

$$V^S(X \cup \{e_1\}) + V^S(X \cup \{e_2\}) = 43/64 + 31/32 = 1.640625$$

486 Therefore, $V^S(X \cup \{e_1, e_2\}) + V^S(X) > V^S(X \cup \{e_1\}) + V^S(X \cup \{e_2\})$, which concludes the
487 proof.

488 D.3 Proof of Proposition 2.3

489 For the proof of Proposition 2.3 we make one assumption about the distribution of edge rejections
490 and failures: querying *additional* edges cannot increase the overall probability of rejection or failure
491 for any edge.

492 **Assumption D.1.** Let $\mathbf{q}, \mathbf{r} \in \{0, 1\}^{|E|}$ denote initial edge queries and responses. Let $\mathbf{q}'$ be additional
493 edges, such that $\mathbf{q} + \mathbf{q}' \in \{0, 1\}^{|E|}$ denotes an augmented edge set; let $\mathbf{r}' \in \{0, 1\}^{|E|}$ denote the
494 responses to edges $\mathbf{q}'$ only. We assume that for any such $\mathbf{q}, \mathbf{r}$, and $\mathbf{q}'$,

$$\mathbb{E}[\mathbf{r} + \mathbf{f} \mid \mathbf{q}, \mathbf{r}] \geq \mathbb{E}[\mathbf{r} + \mathbf{r}' + \mathbf{f} \mid \mathbf{q} + \mathbf{q}', \mathbf{r}].$$

495 Intuitively, Assumption D.1 excludes distributions where queries arbitrarily increase edge failure
496 or rejection. For example, Assumption D.1 disallows the following distribution: suppose all edges
497 are independent; all queried edges are accepted ($P(\mathbf{r}_e = 1 \mid \mathbf{q}) = 0$ for all $\mathbf{q}$), all accepted edges
498 have failure probability 0.5 ($P(\mathbf{f}_e = 1 \mid \mathbf{q}_e = 1, \mathbf{r}_e = 0) = 0.5$), and all non-queried edges have
499 failure probability 0.1 ($P(\mathbf{f}_e = 1 \mid \mathbf{q}_e = \mathbf{r}_e = 0) = 0.1$). In this case, if an edge is not queried, then
500 it has overall rejection or failure probability 0.1 (i.e., $\mathbb{E}[\mathbf{r}_e + \mathbf{f}_e \mid \mathbf{q}, \mathbf{r}] = 0.1$ with $\mathbf{q}_e = 0$); if this
501 edge is queried, then it has rejection or failure probability 0.5 (i.e., $\mathbb{E}[\mathbf{r}_e + \mathbf{r}'_e + \mathbf{f}_e \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] = 0.5$
502 with $\mathbf{q}'_e = 1$).

503 First we prove a handful of useful results.

Definition D.2 (Edge Independence). Two edges $e, e' \in E$ are independent if (a) their rejection
distributions are conditionally independent, given whether or not they were queried:

$$\mathbf{r}_e \perp\!\!\!\perp \mathbf{r}_{e'} \mid \mathbf{q}_e \quad \text{and} \quad \mathbf{r}_e \perp\!\!\!\perp \mathbf{r}_{e'} \mid \mathbf{q}_{e'}$$

and (b) their failure distributions are conditionally independent, given whether or not they were
queried and rejected:

$$\mathbf{f}_e \perp\!\!\!\perp \mathbf{f}_{e'} \mid \mathbf{q}_e, \mathbf{r}_e \quad \text{and} \quad \mathbf{f}_e \perp\!\!\!\perp \mathbf{f}_{e'} \mid \mathbf{q}_{e'}, \mathbf{r}_{e'}.$$

504 **Lemma D.3.** If all edges are independent, then additional edge queries cannot decrease expected
505 post-match cycle and chain weights. Formally,

$$\mathbb{E}[F(c, \mathbf{r} + \mathbf{f}) \mid \mathbf{q}, \mathbf{r}] \leq \mathbb{E}[F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \mid \mathbf{q} + \mathbf{q}', \mathbf{r}]$$

506 for any $\mathbf{q}, \mathbf{q}' \in \{0, 1\}^{|E|}$ such that $\mathbf{q} + \mathbf{q}' \in \{0, 1\}^{|E|}$, for any $\mathbf{r} \in \{0, 1\}^{|E|}$, and for all $c \in \mathcal{C}$.

507 *Proof.* We address cycles and chains separately.

508 **Cycles.** Conditional on fixed $\mathbf{q}$ and $\mathbf{r}$, the expected weight of cycle $c = (e_1, \dots, e_L)$ is expressed
 509 as

$$\begin{aligned}\mathbb{E}[F(c, \mathbf{r} + \mathbf{f}) \mid \mathbf{q}, \mathbf{r}] &= \left(\sum_{e \in c} w_e \right) \mathbb{E} \left[\prod_{e \in c} (1 - \mathbf{r}_e - \mathbf{f}_e) \mid \mathbf{q}, \mathbf{r} \right] \\ &= \left(\sum_{e \in c} w_e \right) \prod_{e \in c} (1 - \mathbb{E}[\mathbf{r}_e + \mathbf{f}_e \mid \mathbf{q}, \mathbf{r}])\end{aligned}$$

510 where the second step is due to the fact that all $\mathbf{f}_e$ are independent. Similarly, for fixed $\mathbf{q}'$,

$$\mathbb{E}[F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] = \left(\sum_{e \in c} w_e \right) \prod_{e \in c} (1 - \mathbb{E}[\mathbf{r}_e + \mathbf{r}'_e + \mathbf{f}_e \mid \mathbf{q} + \mathbf{q}', \mathbf{r}]) .$$

Due to Assumption D.1, the following inequality holds for all edges $e \in E$

$$\mathbb{E}[\mathbf{r}_e + \mathbf{f}_e \mid \mathbf{q}, \mathbf{r}] \geq \mathbb{E}[\mathbf{r}_e + \mathbf{r}'_e + \mathbf{f}_e \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] ,$$

and it follows that

$$\mathbb{E}[F(c, \mathbf{r} + \mathbf{f}) \mid \mathbf{q}, \mathbf{r}] \leq \mathbb{E}[F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] .$$

511 **Chains.** Similarly, the expected weight of chain $c = (e_1, \dots, e_L)$ is expressed as

$$\begin{aligned}\mathbb{E}[F(c, \mathbf{r} + \mathbf{f}) \mid \mathbf{q}, \mathbf{r}] &= \sum_{k=1}^L \left(\sum_{j=1}^k w_j \right) \mathbb{E} \left[\prod_{j=1}^k (1 - \mathbf{r}_{e_j} - \mathbf{f}_{e_j}) \mid \mathbf{q}, \mathbf{r} \right] \\ &= \sum_{k=1}^L \left(\sum_{j=1}^k w_j \right) \prod_{j=1}^k (1 - \mathbb{E}[\mathbf{r}_{e_j} + \mathbf{f}_{e_j} \mid \mathbf{q}, \mathbf{r}]) ,\end{aligned}$$

512 where the second step is due to the fact that $\mathbf{f}_e$ are independent. Similarly,

$$\mathbb{E}[F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] = \sum_{k=1}^L \left(\sum_{j=1}^k w_j \right) \prod_{j=1}^k (1 - \mathbb{E}[\mathbf{r}_{e_j} + \mathbf{r}'_{e_j} + \mathbf{f}_{e_j} \mid \mathbf{q} + \mathbf{q}', \mathbf{r}]) .$$

as before, due to Assumption D.1 it follows that

$$\mathbb{E}[F(c, \mathbf{r} + \mathbf{f}) \mid \mathbf{q}, \mathbf{r}] \leq \mathbb{E}[F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \mid \mathbf{q} + \mathbf{q}', \mathbf{r}] .$$

513

□

Lemma D.4. *With a failure-aware matching policy, and if all edges are independent, adding a single edge to any edge query set weakly improves the objective of Problem 1. Formally, for any $\mathbf{q}, \mathbf{q}' \in \{0, 1\}^{|E|}$ with $\mathbf{q} + \mathbf{q}' \in \{0, 1\}^{|E|}$ and $|\mathbf{q}'| = 1$, and $M(\mathbf{r}) \equiv M^{\text{FA}}(\mathbf{r})$,*

$$V^S(\mathbf{q}) \leq V^S(\mathbf{q} + \mathbf{q}')$$

514 *Proof.* The objective of Problem 1 for edge set $\mathbf{q}$ is expressed as

$$\begin{aligned}V^S(\mathbf{q}) &= \mathbb{E}_{\mathbf{r} \mid \mathbf{q}} \left[\mathbb{E}_{\mathbf{f} \mid \mathbf{q}, \mathbf{r}} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r}) F(c, \mathbf{r} + \mathbf{f}) \right] \right] \\ &= \sum_{\mathbf{r} \in \{0, 1\}^{|\mathbf{q}|}} P_{\mathbf{q}}(\mathbf{r}) \mathbb{E}_{\mathbf{f} \mid \mathbf{q}, \mathbf{r}} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r}) F(c, \mathbf{r} + \mathbf{f}) \right] \\ &= \sum_{\mathbf{r} \in \{0, 1\}^{|\mathbf{q}|}} P_{\mathbf{q}}(\mathbf{r}) \sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r}) \mathbb{E}_{\mathbf{f} \mid \mathbf{q}, \mathbf{r}} [F(c, \mathbf{r} + \mathbf{f})]\end{aligned}$$

515 For edge set $\mathbf{q} + \mathbf{q}'$ we partition response variables into $\mathbf{r}, \mathbf{r}' \in \{0, 1\}^{|E|}$, where $\mathbf{r}_e$ is the response
 516 variable for all edges $e \in \mathbf{q}$, and $\mathbf{r}_e = 0$ for all other edges (including the edge in $\mathbf{q}'$). Similarly, $\mathbf{r}'_e$ is

517 the response variable for edge q' , and $r'_e = 0$ for all other edges. The objective of $q + q'$ is expressed
 518 as

$$\begin{aligned} V^S(q + q') &= \mathbb{E}_{\mathbf{r}, \mathbf{r}' | q + q'} \left[\mathbb{E}_{\mathbf{f} | q + q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r} + \mathbf{r}') F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \right] \\ &= \sum_{\mathbf{r} \in \{0,1\}^{|q|}} P_{q+q'}(\mathbf{r}) \mathbb{E}_{\mathbf{r}' | q+q'} \left[\mathbb{E}_{\mathbf{f} | q+q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r} + \mathbf{r}')^\top F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \right] \\ &= \sum_{\mathbf{r} \in \{0,1\}^{|q|}} P_q(\mathbf{r}) \mathbb{E}_{\mathbf{r}' | q+q'} \left[\mathbb{E}_{\mathbf{f} | q+q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r} + \mathbf{r}') F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \right], \end{aligned}$$

519 where in the final line we replace $P_{q+q'}(\mathbf{r})$ with $P_q(\mathbf{r})$, because each r_e is conditionally independent,
 520 given q_e .

Next, by definition

$$\mathbb{E}_{\mathbf{f} | q + q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r} + \mathbf{r}') F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \geq \mathbb{E}_{\mathbf{f} | q + q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} x_c F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \quad \forall \mathbf{x} \in \mathcal{M}.$$

521 That is, M^{FA} is guaranteed to maximize this expectation, and thus

$$V^S(q + q') \geq \sum_{\mathbf{r} \in \{0,1\}^{|q|}} P_q(\mathbf{r}) \mathbb{E}_{\mathbf{r}' | q+q'} \left[\mathbb{E}_{\mathbf{f} | q+q', \mathbf{r} + \mathbf{r}'} \left[\sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r}) F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f}) \right] \right] \quad (\text{B})$$

$$= \sum_{\mathbf{r} \in \{0,1\}^{|q|}} P_q(\mathbf{r}) \sum_{c \in \mathcal{C}} M_c^{\text{FA}}(\mathbf{r}) \mathbb{E}_{\mathbf{r}' | q+q'} \left[\mathbb{E}_{\mathbf{f} | q+q', \mathbf{r} + \mathbf{r}'} [F(c, \mathbf{r} + \mathbf{r}' + \mathbf{f})] \right] \quad (\text{C})$$

Finally, combining (B) and (C) with Lemma D.3, the following inequality holds

$$V^S(q) \leq V^S(q + q').$$

522

□

523 Using the above lemmas, the proof of Proposition 2.3 is straightforward:

524 **Proposition 2.3** *With a failure-aware matching policy, if all edges are independent, the objective*
 525 *of Problem 1 is monotonic in the set of queried edges.*

Proof. Let $q', q'' \in \mathcal{E}$ be two edge sets such that $q' \subseteq q''$. It remains to show that, with matching policy $M(\mathbf{r}) \equiv M^{\text{FA}}(\mathbf{r})$,

$$V^S(q'') \leq V^S(q').$$

526 First note that because $\mathcal{E}$ is a matroid, there is a sequence of edges $(q^{e_1}, \dots, q^{e_L})$ (with each
 527 $|q^{e_i}| = 1$) such that $q'' + q^{e_1} + \dots + q^{e_L} = q'$. Due to Lemma D.4, the following sequence of
 528 inequalities hold:

$$\begin{aligned} V(q'') &\leq V(q'' + q^{e_1}) \\ &\leq V(q'' + q^{e_1} + q^{e_2}) \\ &\dots \\ &\leq V(q'' + q^{e_1} + \dots + q^{e_L}) \\ &= V(q') \end{aligned}$$

529 which concludes the proof. □

530 E Algorithm Descriptions

531 Here we describe more explicitly the algorithms for Greedy and MCTS, for both the single-stage
 532 and multi-stage settings.

ALGORITHM 3: Greedy: Greedy Search Heuristic for Single-Stage Edge Selection

(input) $\mathcal{E}$: legal edge sets

```
 $q^R \leftarrow \mathbf{0}$    the root node (no edges)
 $V^* \leftarrow$  objective value of  $q^R$  Problem 1
while  $q^R$  has children do
     $q' \leftarrow$  child node of  $q^R$  with maximal objective value in Problem 1
     $q^R \leftarrow q'$ 
return  $q^R$ 
```

E.1 Algorithms for Edge Selection

Algorithm 3 gives a pseudocode description of Greedy for the single-stage setting.

E.2 Multi-Stage Edge Selection

Multi-Stage MCTS. The multi-stage version of MCTS differs from the single-stage version in that each node of the search tree corresponds to both a set of queried edges and a set of observed rejections.

ALGORITHM 4: Multi-Stage MCTS: Tree Search Heuristic for Multi-Stage Edge Selection

TODO: this has not been updated yet!!! (input)
 $\mathcal{E}$: legal edge sets
(input) K : maximum size of any legal edge set
(input) T : time limit per level
(input) L : number of look-ahead levels

```
 $q^R \leftarrow \mathbf{0}$    root node (no edges)
 $q^* \leftarrow \mathbf{0}$    the best visited node
 $V^* \leftarrow$  objective value of  $q^*$ 
for  $N = 1, \dots, K$  do
     $M \leftarrow \min\{N + L, K\}$ 
     $Q \leftarrow$  all nodes in levels  $N$  to  $M$ 
     $U[q] \leftarrow 0 \forall q \in Q$    UCB value estimate
     $V[q] \leftarrow 0 \forall q \in Q$    objective value
     $N[q] \leftarrow 0 \forall q \in Q$    number of visits
    while less than time  $T$  has passed do
         $\text{Sample}(q^R, M)$ 
         $q^R \leftarrow \arg \max_{q \in C(q^R)} U[q]$ 
        Delete  $U[\cdot]$ ,  $V[\cdot]$ , and  $N[\cdot]$ 
return  $q^*$ 
```

ALGORITHM 5: MultiSample: Sampling function used by Multi-Stage MCTS

TODO: this has not been updated yet!!!

```
(input)  $q$ : tree node
(input)  $M$ : maximum level to sample from
 $N[q] \leftarrow N[q] + 1$ 
 $V[q] \leftarrow$  objective of edge set  $q$  in Problem 1
if  $V[q] > V^*$  then
     $q^* \leftarrow q$ ,  $V^* \leftarrow V[q]$ 
if  $q$  has no children then
    return  $V[q]$ 
if  $q$  has children then
    if  $|q| < M$  then
         $q' \leftarrow \arg \max_{q \in C(q^R)} U[q]$ 
         $U[q] \leftarrow U[q] + \text{Sample}(q', M)$ 
    else
         $q' \leftarrow$  a random descendent of  $q$  at any level
         $V' \leftarrow$  objective value of  $q'$  in Problem 1
        if  $V' > V^*$  then
             $q^* \leftarrow q'$ ,  $V^* \leftarrow V'$ 
         $U[q] \leftarrow U[q] + V'$ 
```

Multi-Stage Greedy. Algorithm 6 gives a pseudocode description of the multi-stage version of Greedy. This search heuristic returns the next edge to query with the highest expected final matching weight, *ignoring all future queries*. In other words, this approach treats every edge as the *last* edge; one might call this heuristic “myopic” as well as greedy.

ALGORITHM 6: Greedy Heuristic for Multi-Stage Edge Selection

(input) $\mathcal{E}$: legal edge sets

(input) q : previously-queried edges

(input) r : previously-observed rejections

$e^* \leftarrow \emptyset$ $V^* \leftarrow 0$

for all q' *in* q 's children **do**

$e' \leftarrow$ the new edge queried in child node q' $r^A \leftarrow r$

$r^R \leftarrow r$

$r_e^A \leftarrow 1$ (response scenario where e' is accepted)

$r_e^R \leftarrow 1$ (response scenario where e' is rejected)

$p^A \leftarrow$ probability that e is accepted, conditional on r

$p^R \leftarrow$ probability that e is accepted, conditional on r

$V' \leftarrow p^A \cdot W(M(r^A); q', r^A) + p^R W(M(r^R); q', q^R)$ (value of querying edge e')

if $V' > V^*$ **then**

$e^* \leftarrow e'$

$V^* \leftarrow V'$

return e^*
