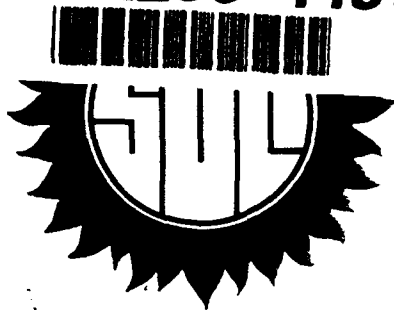


(27)

AD-A236 419.



Systems
Optimization
Laboratory

DTIC

S

ELECTR
JUN 04 1991

D

c



Accession For	
DTIC GRA&I	☒
DTIC L&S	☐
Unannounced	☐
Justification	
By	
Distribution	
Availability Codes	
Avail and/or	
Dist	Special
A-1	

Approved for Release
Distribution Unlimited

Department of Operations Research
Stanford University
Stanford, CA 94305

91-00673



91 5 29 034

**SYSTEMS OPTIMIZATION LABORATORY
DEPARTMENT OF OPERATIONS RESEARCH
STANFORD UNIVERSITY
STANFORD, CALIFORNIA 94305-4022**

**Monte Carlo (Importance) Sampling within a Benders
Decomposition Algorithm for Stochastic Linear Programs
Extended Version: Including Results of Large-Scale Problems**

by
Gerd Infanger

TECHNICAL REPORT SOL 91-6

March 1991

Research and reproduction of this report were partially supported by the Office of Naval Research Contract N00014-89-J-1659; the National Science Foundation Grants ECS-8906260, DMS-8913089; the Electric Power Research Institute Contract RP 8010-09 at Stanford University. Additional research was supported by the Electric Power Research Institute Contract CSA-4O05335, and the Austrian Science Foundation, "Fonds zur Förderung der wissenschaftlichen Forschung," Grant J0323-PHY.

Any opinions, findings, and conclusions or recommendations expressed in this publication are those of the author and do NOT necessarily reflect the views of the above sponsors.

Reproduction in whole or in part is permitted for any purposes of the United States Government. This document has been approved for public release and sale; its distribution is unlimited.

**Monte Carlo (Importance) Sampling within a Benders
Decomposition Algorithm for Stochastic Linear Programs
Extended Version: Including Results of Large-Scale Problems ***

Gerd Infanger

Department of Operations Research, Stanford University

March 1991

Abstract

The paper focuses on Benders decomposition techniques and Monte Carlo sampling (importance sampling) for solving two-stage stochastic linear programs with recourse, a method first introduced by George B. Dantzig and Peter Glynn (1990). The algorithm is discussed and further developed. The paper gives a complete presentation of the method as it is currently implemented. Numerical results from test problems of different areas are presented. Using small test problems we compare the solutions obtained by the algorithm with the universe solutions. We present the solution of large-scale problems with numerous stochastic parameters which in the deterministic equivalent formulation would have billions of constraints. The problems concern expansion planning of electric utilities with uncertainty in the availabilities of generators and transmission lines and portfolio management with uncertainty in the future returns.

* This paper is an update of Technical Report SOL 89-13R, August 1990, Department of Operations Research, Stanford University, with extended numerical results including numerical results of large-scale test problems.

1. Introduction

A stochastic linear program is a linear program whose parameters (coefficients, right hand sides) are uncertain. The uncertain parameters are assumed to be known only by their distributions. That means that the values of some functions are numerical characteristics of random phenomena, e.g. mathematical expectations of functions dependent on decision variables and random parameters.

Suppose a function $z = E C(V)$ is an expectation of a function $C(v^\omega), \omega \in \Omega$. V is a random parameter which has outcomes v^ω . Ω is the set of all possible random events. It can be finite, infinite, discrete or continuous. In the continuous case the computation of the expected value requires to solve the integral:

$$E C(V) = \int C(v^\omega) P(d\omega)$$

with P being the probability measure.

In a general case V would consist of several components, e.g. $V = (V_1, \dots, V_h)$ with outcomes v^ω which we also will denote only by lower case letters, e.g. $v = (v_1, \dots, v_h)$ and $p(v^\omega)$ alias $p(v)$ would denote the corresponding density function. We assume the components of V to be independent. In addition we will construct Ω by crossing the sets of outcomes Ω_i for vector entry $v_i, i = 1, \dots, h$ as

$$\Omega = \Omega_1 \times \Omega_2 \times \dots \times \Omega_h.$$

In this case the above mentioned integral takes the form of a multiple integral:

$$E C(V) = \int \dots \int C(v) p(v) dv_1 \dots dv_h$$

In the case of Ω being discrete and finite the expectation can be computed by a multiple sum:

$$E C(V) = \sum_{v_1} \dots \sum_{v_h} C(v) p(v).$$

The main difficulties in stochastic linear programming deal with the evaluation of the multiple integral or the multiple sum. The numerical computation of the expectation requires a large number of function evaluations and each function evaluation means a linear program to be solved. Different approaches attack this problem, e.g. Birge (1985), Birge and Wets (1986), Birge and Wallace (1988), Frauendorfer (1988), Frauendorfer and Kall (1988), Ermoliev (1983), Hige and Sen (1989), Kall (1979), Pereira et al. (1989), Rockafellar and Wets (1989), Ruszczynski (1986), Wets (1984), and others. See Ermoliev and Wets (1988) for references. We follow the concept of Dantzig et al. (1989) and Dantzig and Glynn (1990).

2. Two Stage Stochastic Linear Program

An important class of models concerns dynamic linear programs. Variables which describe activities initiated at time t have coefficients at time t and $t + 1$. Deterministic dynamic linear programs appear as staircase problems. The simplest staircase problem is that with two stages: X denotes the first, Y the second stage decision variables, A, b represent the coefficients and right hand sides of the first stage constraints and D, d concern the second period constraints together with B which couples the two periods. c, f are the objective function coefficients.

In the deterministic case c, f, A, b, B, D, d are known with certainty to the planner. In the stochastic case, the parameters of the second stage are not known to the planner at time $t = 1$, but will be known at time $t = 2$. At time $t = 1$ only the distributions of these parameters are assumed to be known. The second stage parameters can be seen as random variables which get certain outcomes with certain probabilities. We denote a certain outcome of these parameters with ω and the corresponding probability with $p^\omega, \omega \in \Omega$, the set of possible outcomes.

$$\begin{array}{ll}
\min Z = cX + E^\omega(fY^\omega) & \\
s/t & \\
\quad AX & = b \\
\quad - B^\omega X + DY^\omega & = d^\omega \\
\quad X, Y^\omega & \geq 0, \quad \omega \in \Omega
\end{array} \quad (1)$$

In (1) a two stage staircase problem is transformed into a two stage stochastic linear program and the parameters d and B being random variables. Given the two stage stochastic linear program one wants to make a decision X which is feasible for all scenarios and has the minimum expected costs.

We consider the case of Ω being discrete and finite, e.g. $\Omega = (1, \dots, K)$, the parameter ω takes on K values. Then we can formulate an equivalent deterministic problem to the stochastic linear problem. This is tractable if K is small. For K scenarios the deterministic equivalent problem is given in (2).

$$\begin{array}{ll}
\min Z = cX + p^1 fY^1 + p^2 fY^2 + \dots + p^K fY^K & \\
s/t & \\
\quad AX & = b \\
\quad -B^1 X + DY^1 & = d^1 \\
\quad -B^2 X & + DY^2 = d^2 \\
\quad \vdots & \ddots \vdots \\
\quad -B^K X & + DY^K = d^K \\
\quad X, Y^1, Y^2, \dots, Y^K & \geq 0
\end{array} \quad (2)$$

Two stage stochastic linear programs were first studied in Dantzig (1955) and then developed by many authors. The method which we want to apply here is using Benders (1962) decomposition. Van Slyke and Wets (1969) suggested to express the impact of the second period by a scalar θ and "cuts", which are necessary conditions to the problem and are expressed only in terms of the first period variables X and θ . Benders decomposition splits the original problem into a master problem and a subproblem which decomposes into a series of independent subproblems, one for to each $\omega \in \Omega$. According to the L-shaped method the master problem, the sub problems and the cuts are represented in (3), (4) and (5).

The master problem:

$$\begin{aligned} \min z_M &= cX + \theta \\ \text{s/t} \quad &AX = b \\ &-G^\ell X + \alpha^\ell \theta \geq g^\ell, \ell = 1, \dots, L \\ &X, \theta \geq 0 \end{aligned} \quad (3)$$

The sub problems:

$$\begin{aligned} \min z^\omega &= p^\omega f Y^\omega \\ \text{s/t} \quad p^\omega \pi^\omega : &DY^\omega = d^\omega + B^\omega X \\ &Y^\omega \geq 0, \omega \in \Omega, \text{ e.g. } \Omega = \{1, 2, \dots, K\} \end{aligned} \quad (4)$$

where $p^\omega \pi^{\omega*}$ is the optimal dual solution of subproblem ω .

The cuts:

$$\begin{aligned} g &= \sum_{\omega} p^\omega \pi^{\omega*} d^\omega = E(\pi^{\omega*} d^\omega) \\ G &= \sum_{\omega} p^\omega \pi^{\omega*} B^\omega = E(\pi^{\omega*} B^\omega) \end{aligned} \quad (5)$$

$$\alpha^\ell = 0 \dots \text{feasibility cut}$$

$$\alpha^\ell = 1 \dots \text{optimality cut}$$

By solving the master problem we obtain a solution X . Given X we can solve K subproblems $\omega \in \Omega$ to compute a cut. The cut is a lower bound on the expected value of the second stage costs represented as a function of X . Cuts are sequentially added to the master problem and new values of X are obtained until the optimality criterion is met. We distinguish between two types of cuts, feasibility cuts and optimality cuts. The first refers to infeasible subproblems for a given X and the latter to feasible and optimum subproblems, given X .

If the expected values z , G , and g are computed exactly, that is, by evaluating all scenarios $\omega \in \Omega$, we refer to it as the universe case. As we will see later the number of scenarios easily gets out of hand and it is not always possible to solve the universe case. Therefore methods are sought which guarantee a satisfying solution without solving the universe case.

3. Monte Carlo Sampling

Each iteration of Benders decomposition requires the computation of expected values, e.g. the subproblem costs, the coefficients and right hand sides of the cuts. For each outcome $\omega \in \Omega$ a linear program has to be solved. The expected value of the subproblem costs is denoted by

$$z = E C(v^\omega) = E fY^{*\omega}, \quad \omega \in \Omega,$$

with $Y^{*\omega}$ being the optimum solution of subproblem ω . The number of elements of Ω is determined by the dimensionality of the stochastic vector $V = (V_1, \dots, V_h)$. Typically the dimension h of V is quite large. For example, in expansion planning problems of electric power systems one component of V denotes the availability of one type of generators or one demand of power in a node of a multi-area supply network or the availability of one type of transmission line connecting two nodes. Consider several nodes and arcs and one demand and some options of generators at each node. The number of scenarios K in the universe case quickly gets out of hand, even if the distribution of each component of V is determined by just a small number K^i of discrete points. Suppose e.g. $h = 20$ and $K^i = 5$, $i = 1, \dots, 20$. Then the total number of terms in the expected value calculations is $K = 5^{20} \approx 10^{14}$, which is not practically solvable by direct summation. Monte Carlo methods appear promising to compute multiple integrals or multiple sums for h large (Davis and Rabinowitz (1984)). See Hammersly and Handscomb (1964) for a description of Monte Carlo sampling techniques.

3.1 Crude Monte Carlo

Suppose $v^\omega, \omega = 1, \dots, n$ are scenarios, sampled independently from their joint probability mass function, then $C^\omega = C(v^\omega)$ are independent random variates with expectation z .

$$\bar{z} = (1/n) \sum_{\omega=1}^n C^\omega \tag{6}$$

is an unbiased estimator of z and its variance is

$$\sigma_{\bar{z}}^2 = \sigma^2/n$$

$$\sigma^2 = \text{var}(C(V)).$$

Thus the standard error is decreasing with sample size n by $n^{-0.5}$. The convergence rate of $\bar{z}$ to z is independent of the dimension h of the random vector V .

3.2 Importance Sampling

We rewrite

$$z = \sum_{\omega \in \Omega} C(v^\omega) p(v^\omega) = \sum_{\omega \in \Omega} \frac{C(v^\omega) p(v^\omega) q(v^\omega)}{q(v^\omega)}$$

by introducing a probability mass function $q(v^\omega)$. We can see q as a probability mass function of a random vector W , therefore by change of variables

$$z = E \frac{C(W) p(W)}{q(W)}.$$

We obtain a new estimator of z ,

$$\bar{z} = \frac{1}{n} \sum_{\omega=1}^n \frac{C(w^\omega) p(w^\omega)}{q(w^\omega)}$$

which has a variance of

$$\text{var}(\bar{z}) = \frac{1}{n} \sum_{\omega \in \Omega} \left(\frac{C(w^\omega) p(w^\omega)}{q(w^\omega)} - z \right)^2 q(w^\omega).$$

Choosing $q^*(w^\omega) = \frac{C(w^\omega) p(w^\omega)}{\sum_{\omega \in \Omega} C(w^\omega) p(w^\omega)}$ would lead to $\text{var}(\bar{z}) = 0$, that means we could get a perfect estimate of the multiple sum just by one single observation. However this is practically useless, since to sample $C.p/q$ we have to know q and to determine q we need to know $z = \sum_{\omega \in \Omega} C(w^\omega) p(w^\omega)$, which we eventually want to compute. Nevertheless this result helps to derive some heuristics of how to choose q : It should be approximately proportional to the product $C(w^\omega) p(w^\omega)$ and have a

form which can be integrated analytically. For instance using the additive (separable in the components of the stochastic vector) approximation:

$$C(V) \approx \sum_{i=1}^h C_i(V_i)$$

could be a possible way to compute a proper q :

$$q(w^\omega) \approx \frac{C(w^\omega)p(w^\omega)}{\sum_{i=1}^h \sum_{\omega \in \Omega_i} C_i(w^\omega)}.$$

In this case one has to solve only h 1-dimensional sums instead of 1 h -dimensional sum. Depending on how well the additive model approximates the original cost surface the above mentioned estimator will lead to smaller variances compared to crude Monte Carlo sampling. Of course if the original cost surface has the property of additivity (separability) no sampling is required, as the multiple sum is computed exactly by h 1-dimensional sums.

The advantage of this approach lies in the fact that even if the additive model is a bad approximation to the cost surface the method works. The price that has to be paid is a high sample size. The variance reduction compared to crude Monte Carlo will be small. For the theory of importance sampling we refer to Glynn and Iglehart (1989). See also Dantzig and Glynn (1990).

R. Entriken and M. Nakayama in Dantzig et. al. (1989) developed an importance sampling scheme using an additive model to approximate the cost function $E C(V)$. Actually $C(v)$ is approximated by a marginal cost model, considering marginal costs in each dimension i of V and a base case, the point from which the approximation is developed. We will use this approach here. As we employ importance sampling within the Benders decomposition algorithm the costs $C(v, \hat{X})$, the approximation of the costs $\Gamma(v, \hat{X})$ and thus the importance distribution of $q(v, \hat{X})$ depend also on $\hat{X}$, the current solution of the master problem. Introducing the costs of the base case $C(\tau, \hat{X})$ makes the model more sensitive to the impact of the

stochastic variables V .

$$C(V, \hat{X}) \approx \Gamma(V, \hat{X}) = C(\tau, \hat{X}) + \sum_{i=1}^h M_i(V_i, \hat{X}), \quad (7)$$

$$M_i(V_i, \hat{X}) = C(\tau_1, \dots, \tau_{i-1}, V_i, \tau_{i+1}, \dots, \tau_h, \hat{X}) - C(\tau, \hat{X}).$$

$\tau = (\tau_1, \dots, \tau_h)$ can be any arbitrary chosen point out of the set of values v_i , $i = 1, \dots, h$. For example we choose τ_i as that outcome of V_i which leads to the lowest costs, *ceteris paribus*. These values can be found easily. Note that the second stage costs are computed by a linear program, where the uncertain parameters appear in right hand side. Therefore the second stage costs are convex in the stochastic parameters V . The sign of the dual variables associated with the stochastic parameter indicate the direction to lowest costs. In the context of expansion planning of power systems this means selecting respectively lowest demands and highest availabilities of generators and transmission lines.

Defining

$$\bar{M}_i(\hat{X}) = E M_i(V_i, \hat{X}) = \sum_{\omega \in \Omega_i} M_i(v_i^\omega, \hat{X}) p(v_i^\omega) \quad (8)$$

and

$$F(v^\omega, \hat{X}) = \frac{C(v^\omega, \hat{X}) - C(\tau, \hat{X})}{\sum_{i=1}^h M_i(v_i^\omega, \hat{X})} \quad (9)$$

where we assume that

$$\sum_{i=1}^h M_i(v_i^\omega, \hat{X}) > 0. \text{ that means at least one } M_i(v_i^\omega, \hat{X}) > 0,$$

we can express the expected value of the costs in the following form:

$$z(\hat{X}) = C(\tau, \hat{X}) + \sum_{i=1}^h \bar{M}_i(\hat{X}) \sum_{\omega \in \Omega} F(v^\omega, \hat{X}) \frac{M_i(v_i^\omega, \hat{X})}{\bar{M}_i(\hat{X})} \prod_{j=1}^h p_j(v_j^\omega). \quad (10)$$

Note that this formulation consists of a constant term and a sum of h expectations. Given a fixed sample size n we partition n into h sub-samples, with sample sizes n_i , $i = 1, \dots, h$ such that $\sum n_i = n$ and $n_i \geq 1$, $i = 1, \dots, h$ and n_i being approximately proportional to $\bar{M}_i$. The h expectations are separately approximated

by sampling using marginal densities. The i -th expectation corresponds of course to the i -th component of V . Generating sample points in the i -th expectation we use the importance density $(p_i M_i / \bar{M}_i)$ for sampling the i -th components of V and the original marginal densities for any other components. Denoting

$$\mu_i(\hat{X}) = \frac{1}{n_i} \sum_{j=1}^{n_i} F(v^j, \hat{X}) \quad (11)$$

the estimate of the i -th sum, we obtain

$$\bar{z}(\hat{X}) = C(\tau, \hat{X}) + \sum_{i=1}^h \bar{M}_i(\hat{X}) \mu_i(\hat{X}), \quad (12)$$

the estimated expected value of the second stage costs $\bar{z}(\hat{X})$.

Let $\bar{\sigma}_i^2(\hat{X})$ be the estimated sample variance of the i -th expectation, where $\sigma_i^2(\hat{X}) = 0$ if $n_i = 1$. The estimated variance of the mean, $\sigma_z^2(\hat{X})$, is then given by

$$\sigma_z^2(\hat{X}) = \sum_{i=1}^h \frac{\bar{M}_i^2(\hat{X}) \bar{\sigma}_i^2(\hat{X})}{n_i}. \quad (13)$$

Using importance sampling one can achieve significant variance reduction. The experiment of *M. Nakayama* in *Dantzig et al. (1989)* claims a variance reduction of 1:20000 using importance sampling versus crude Monte Carlo sampling: For a given and optimal $\hat{X}$ the second stage costs of a multi area expansion planning model with 192 universe scenarios were sampled with a sample size of 10 using both methods and the results compared.

The derivation above concerned the estimation of the expected second stage costs $z(\hat{X})$. To derive a cut we use the same framework analogously. Note that a cut is defined as an outer linearization of the second stage costs represented as a function of X , the first stage variables. At $\hat{X}$, the value of the cut is exactly the expected second stage costs $z(\hat{X})$. Note also that any choice of q is a valid choice. As we do not want to derive different importance distributions for the coefficients and the right hand side of a cut, we use the q already at hand from the

cost estimation. Therefore we employ directly the cost approximation scheme and the importance distribution to compute the gradient and the right hand side of a cut. With $B(v^\omega) := B^\omega$ and $d(v^\omega) := d^\omega$ being the outcome of B and d in scenarios $\omega, \omega \in \Omega$ and $\pi^*(v^\omega, \hat{X}) := \pi^{\omega*}(\hat{X})$, the optimum dual solution in scenario ω , we define

$$F^G(v^\omega, \hat{X}) = \frac{\pi^*(v^\omega, \hat{X})B(v^\omega) - \pi^*(\tau, \hat{X})B(\tau)}{\sum_{i=1}^h M_i(v_i^\omega, \hat{X})} \quad (14)$$

$$F^g(v^\omega, \hat{X}) = \frac{\pi^*(v^\omega, \hat{X})d(v^\omega) - \pi^*(\tau, \hat{X})d(\tau)}{\sum_{i=1}^h M_i(v_i^\omega, \hat{X})} \quad (15)$$

and compute:

$$G(\hat{X}) = \pi^*(\tau, \hat{X})B(\tau) + \sum_{i=1}^h \bar{M}_i(\hat{X}) \sum_{\omega \in \Omega} F^G(v^\omega, \hat{X}) \frac{M_i(v_i^\omega, \hat{X})}{\bar{M}_i(\hat{X})} \prod_{j=1}^h p_j(v_j^\omega) \quad (16)$$

$$g(\hat{X}) = \pi^*(\tau, \hat{X})B(\tau) + \sum_{i=1}^h \bar{M}_i(\hat{X}) \sum_{\omega \in \Omega} F^g(v^\omega, \hat{X}) \frac{M_i(v_i^\omega, \hat{X})}{\bar{M}_i(\hat{X})} \prod_{j=1}^h p_j(v_j^\omega), \quad (17)$$

the coefficients and the right hand side of a cut. We estimate the expected values again by sampling using the same sample points as at hand from the computation of z .

Using Monte Carlo sampling we obtain $\bar{z}(\hat{X}), \bar{G}(\hat{X}), \bar{g}(\hat{X})$, which are approximations of the expected values $z(\hat{X}), G(\hat{X}), g(\hat{X})$. We also obtain the estimated variance of the mean of the second stage costs $\sigma_{\bar{z}}(\hat{X})$. The impact of using approximations instead of the exact parameters on the Benders decomposition algorithm is analyzed in the following section.

4. Benders Decomposition

In the following we will derive the main steps of Benders decomposition algorithm for two stage stochastic linear programs considering the “universe” case, which gives the exact solution of the equivalent deterministic problem (“certainty

equivalent"). We will then analyze the impact of sampling of subproblems on Benders decomposition. See Geoffrion (1970) for a derivation of Benders decomposition algorithm.

Given the equivalent deterministic problem in (2) and assuming K scenarios describe the universe case, we rewrite the problem applying projection onto the X variables and obtain (18). We assume for simplicity that (2) is feasible and has a finite optimum solution.

$$\begin{array}{llll}
 \min Z = & cX & + \text{Min}[p^1 fY^1 + p^2 fY^2 + \dots + p^K fY^K] & \\
 AX = b & & DY^1 & = d^1 + B^1 X \\
 X \geq 0 & & DY^2 & = d^2 + B^2 X \\
 & & \ddots & \vdots \\
 & & & DY^K = d^K + B^K X \\
 & Y^1, Y^2, \dots, & & Y^K \geq 0
 \end{array} \quad (18)$$

The infimal value function in (18) corresponds to the following primal linear problem (19):

$$\begin{array}{llll}
 \min z_P = p^1 fY^1 + p^2 fY^2 + \dots + p^K fY^K = E^\omega(fY^\omega) & & & \\
 p^1 \pi^1 : & DY^1 & & = d^1 + B^1 X \\
 p^2 \pi^2 : & & DY^2 & = d^2 + B^2 X \\
 \vdots & & \ddots & \vdots \\
 p^K \pi^K : & & & DY^K = d^K + B^K X \\
 & Y^1, Y^2, \dots, Y^K \geq 0 & &
 \end{array} \quad (19)$$

and to the dual linear problem (20):

$$\begin{array}{llll}
 \max z_D = & p^1 \pi^1 (d^1 + B^1 X) + p^2 \pi^2 (d^2 + B^2 X) + \dots + p^K \pi^K (d^K + B^K X) & & \\
 \pi^1 D & & & \leq f \\
 & \pi^2 D & & \leq f \\
 & & \ddots & \vdots \\
 & & \pi^K D & \leq f
 \end{array} \quad (20)$$

The primal problem is parameterized in the right hand side by X . The assumption (2) being finite implies that (19) is finite for at least one value of X for which $X \geq 0$ and $AX = b$. Applying the Duality Theorem of Linear Programming

we state that (20) has to be feasible. The feasibility conditions

$$\pi^\omega D - f \leq 0$$

indicate that the feasible region $\{\pi^\omega | \pi^\omega D - f \leq 0\}$ is independent of X and ω and just repeated for each scenario $\omega \in \Omega$.

The assumption (2) being feasible requires feasibility of the primal problem (19) for at least one X for which $X \geq 0$ and $AX = b$. We define $\pi := (\pi^1, \pi^2, \dots, \pi^K)$ to be the vector of dual variables of problem (20). By the Duality Theorem again (20) has to be finite. Let $\pi^j, j = 1, \dots, p$ be the extreme points and $\pi^j, j = p+1, \dots, p+q$ be representatives of the extreme rays of the feasible region of (20), where $\pi^j := (\pi^{1j}, \pi^{2j}, \dots, \pi^{Kj})$. Problem (20) is finite if and only if

$$\begin{aligned} \pi^{\omega j}(d^\omega + B^\omega X) &\leq 0, \quad j = p+1, \dots, p+q \\ \omega &\in \Omega. \end{aligned} \quad (21)$$

Constraints (21) may be appended to problem (18) to ensure that the problem is bounded.

Next we outer linearize the infimal value function in (18), whose value is exactly:

$$\text{Maximum}_{j=1, \dots, p} \sum_{\omega \in \Omega} p^\omega \pi^{\omega j} (d^\omega + B^\omega X). \quad (22)$$

By expressing the infimal value function by the outer linearized dual problem and using θ as the smallest upper bound the problem can be represented in the following form:

$$\begin{aligned} \min Z = & \quad cX + \theta \\ & AX = b \\ & X \geq 0 \\ \theta \geq & \sum_{\omega \in \Omega} p^\omega \pi^{\omega j} (d^\omega + B^\omega X), \quad j = 1, \dots, p \\ & \pi^j (d^\omega + B^\omega X) \leq 0, \quad j = p+1, \dots, p+q, \quad \omega \in \Omega. \end{aligned} \quad (23)$$

Relaxation is applied to solve problem (23) as we do not want to know all $\pi^j, j = 1, \dots, p+q$ in advance: Given a solution $(\hat{X}, \hat{\theta})$ from the master problem

one solves problem (19) or problem (20), actually by solving the individual problems (4) or the dual problems (24) of these:

$$z^{\omega*}(\hat{X}) = \max z_D^{\omega} = \begin{array}{ll} \pi^{\omega} D & \leq f \\ \pi^{\omega} & \geq 0, \end{array} \quad \omega \in \Omega. \quad (24)$$

We call $\pi^{\omega*}(\hat{X})$ the optimum dual solution vector. If primal infeasibility or dual unboundness is detected, with $\pi^{\omega^{\circ}}(\hat{X})$ denoting the corresponding extreme ray, a feasibility cut

$$\pi^{\omega^{\circ}}(\hat{X}) \cdot (d^{\omega} + B^{\omega} X) \leq 0 \quad (25)$$

is added to the master problem. If all primal problems are feasible or all dual problems bounded an optimality cut:

$$\theta \geq \sum_{\omega \in \Omega} p^{\omega} \pi^{\omega*}(\hat{X}) \cdot (d^{\omega} + b^{\omega} X) \quad (26)$$

is added to the master problem. We call

$$L(X) := \sum_{\omega \in \Omega} p^{\omega} \pi^{\omega*}(\hat{X}) \cdot (d^{\omega} + B^{\omega} X) \quad (27)$$

an outer linearization of the second stage costs, which are defined by

$$z(\hat{X}) := \sum_{\omega \in \Omega} z^{\omega*}(\hat{X}). \quad (28)$$

The relation:

$$L(X) \leq z(X) \quad (29)$$

formulates the main property of the outer linearization. Any cut independent of $\hat{X}$ from which it was originally derived, is a valid cut as long as it does not violate the main property of outer linearization.

Benders decomposition algorithm provides upper and lower bounds to the solution in each iteration.

In the l -th iteration

$$LB^\ell := c\hat{X}^\ell + \hat{\theta}^\ell \quad (30)$$

with $\hat{X}^\ell, \hat{\theta}^\ell$ being the optimum solution of the master problem is defined to be a lower bound and

$$UB^\ell := \min\{UB^{\ell-1}, c\hat{X}^\ell + z(\hat{X}^\ell)\}, \quad UB^0 = \infty \quad (31)$$

with $z(\hat{X}^\ell)$ the second stage costs, to be an upper bound to the solution of the problem. If

$$(UB^\ell - LB^\ell)/LB^\ell \leq TOL \quad (32)$$

where TOL is a given tolerance, the problem is said to be solved with a sufficient accuracy.

4.1 Probabilistic Cuts

Employing Monte Carlo sampling techniques means not to solve all problems $\omega \in \Omega$, but solving problems $\omega \in S$, S being a subset of Ω . Instead of the exact expected values $z(\hat{X})$, $G(\hat{X})$, $g(\hat{X})$ we compute the estimates $\bar{z}(\hat{X})$, $\bar{G}(\hat{X})$, $\bar{g}(\hat{X})$ by importance sampling. We also estimate the error of the estimation of $z(\hat{X})$ by the variance $\text{var}(\bar{z}(\hat{X})) = \sigma_z^2(\hat{X})$. Thus e.g. in the case of the second stage costs the estimation results in an estimated mean with some error distribution. There is good reason to assume the error being normally distributed (Davis and Rabinowitz (1984)). We define $\tilde{z}(\hat{X})$ to be random, normally distributed with mean $\bar{z}(\hat{X})$ and variance $\sigma_z^2(\hat{X})$:

$$\tilde{z}(\hat{X}) := N(\bar{z}(\hat{X}), \sigma_z^2(\hat{X})). \quad (33)$$

A cut obtained by sampling differs in general from a cut computed by solving the universe scenarios. The outer linearizations $L(X) = G(\hat{X})X + g(\hat{X})$, with respect to the universe case and $\bar{L}(X) = \bar{G}(\hat{X})X + \bar{g}(\hat{X})$, with respect to the estimation differ in the gradient and the right hand side. At $X = \hat{X}$, where $L(\hat{X}) = z(\hat{X})$ and

$\bar{L}(\hat{X}) = \bar{z}(\hat{X})$ we substitute the variable θ for $\tilde{z}(\hat{X})$ when defining a cut. By this substitution θ takes on the distribution of $\tilde{z}$, therefore $\theta := N(\bar{\theta}, \sigma_{\tilde{z}}^2)$. This is only true at $X = \hat{X}$. However, we assume this error distribution to be constant with respect to X . That means we see the error mainly concentrated in the right hand side of the cut and we assign the variance $\sigma_{\tilde{z}}(\hat{X})$ also to the right hand side and define

$$\tilde{g}(\hat{X}) := N(\bar{g}(\hat{X}), \sigma_{\tilde{z}}(\hat{X})) \quad (34)$$

to be the random right hand side of the cut, normally distributed with mean $\bar{g}(\hat{X})$ and variance $\sigma_{\tilde{z}}^2(\hat{X})$. We can expect that in the final solution cuts will be binding at an X very close to $\hat{X}$, where they were originally derived. The assumption of a constant error distribution of θ is therefore intuitively plausible. See also Dantzig and Glynn (1990) in this respect. In general S is a sufficiently large subset of Ω so that the variance $\sigma_{\tilde{z}}^2$ is small.

Cuts computed by sampling do not necessarily meet the condition of outer linearization. Violating this condition a cut intersects and separates parts of the feasible region of the second stage problem. A sampled cut is therefore not a valid cut.

4.2 Upper and Lower Bounds

For random second stage costs $\tilde{z}(\hat{X}^\ell)$ and random right hand sides $\tilde{g}^\ell$, $\ell = 1, \dots, L$ the upper and lower bounds of the problem as provided by Benders decomposition algorithm are probabilistic.

The Upper Bounds

$$\tilde{U}B^\ell := c\hat{X}^\ell + \tilde{z}(\hat{X}^\ell), \quad \ell = 1, \dots, L \quad (35)$$

are random parameters, normally distributed with means $\bar{U}B^\ell$ and variances $\sigma_{\tilde{z}}(\hat{X}^\ell)$:

$$\tilde{U}B^\ell := N(\bar{U}B^\ell, \sigma_{\tilde{z}}^2(\hat{X}^\ell)) \quad \ell = 1, \dots, L. \quad (36)$$

We define the lowest upper bound to be the upper bound with the lowest mean

$$\tilde{U}B_{\min}^L : \bar{U}B_{\min}^L := \min_{\ell=1, \dots, L} \{\bar{U}B^\ell\} \quad (37)$$

with corresponding variance $\sigma_{\tilde{U}B_{\min}^L}^2$.

The lower bounds are obtained from the solution of the master problem. To determine the distribution of a lower bound consider the master problem at iteration L :

$$\begin{aligned} \tilde{L}B^L = \tilde{z}_M^{*L} = \min \tilde{z}_M^L = & \quad c \quad X + \theta \\ s/t \quad \rho_1^0 : & \quad A \quad X \quad = b \\ \rho^1 : & \quad -G^1 X + \theta \geq \tilde{g}^1 \\ & \quad \vdots \\ \rho^L : & \quad -G^L X + \theta \geq \tilde{g}^L \\ & \quad X, \quad \theta \geq 0 \end{aligned}$$

where L optimality cuts have been added to the originally relaxed master problem. We do not consider feasibility cuts for the following argument, as they are exact. The vector ρ^0 and the scalars $\rho^\ell, \ell = 1, \dots, L$ denote the dual prices. The right hand sides $\tilde{g}^\ell, \ell = 1, \dots, L$ are independent stochastic parameters, normally distributed. We assume independence as the cuts are generated from independent samples, neglecting the dependency that $\hat{X}^\ell, \ell = 1, \dots, L$ are weakly connected by Benders decomposition algorithm.

With the random parameters $\tilde{g}^\ell, \ell = 1, \dots, L$ in the right hand side also the optimum solution $\tilde{z}_M^*$ will be random. We define the optimum solution of the master problem

$$\tilde{z}_M^{*L} := N(\bar{z}_M^L, \text{var}(\tilde{z}_M^{*L})) \quad (39)$$

to be a random parameter, normally distributed, with mean $\bar{z}_M^L$ and variance $\text{var}(\tilde{z}_M^{*L})$. Hence one could experimentally obtain the distribution of $\tilde{z}_M^{*L}$ by randomly varying the right hand sides according to N samples $j = 1, \dots, N$ drawn

from the normal distributions of $\tilde{g}^\ell, \ell = 1, \dots, L$ and by solving the master problem for all N samples. One could estimate the mean and the variance of the distribution from the samples $j = 1, \dots, N$. As this is a very expensive way to obtain an estimate of the lower bound distribution, we proceed instead in the following way. We have already stated that we choose a sample size $|S|$, such that the variances $\sigma_z^\ell, \ell = 1, \dots, L$ are small. If the variances are small we can assume that for all outcomes of the random right hand sides $\tilde{g}^\ell, \ell = 1, \dots, L$, the optimum solution of the master problem has the same basis. Then we can compute the mean of the lower bound estimate:

$$\begin{aligned} \bar{z}_M^{*L} = \min z_M = & \quad c \quad X + \theta \\ s/t \quad \rho^0 : & \quad A \quad X = b \\ \rho^1 : & \quad -G^1 X + \theta \geq \bar{g}^1 \\ & \quad \vdots \\ \rho^L : & \quad -G^L X + \theta \geq \bar{g}^L \\ & \quad X, \quad \theta \geq 0 \end{aligned} \quad (40)$$

by substituting the means $\bar{g}^\ell, \ell = 1, \dots, L$ for the random parameters $\tilde{g}^\ell, \ell = 1, \dots, L$, and the variance $\text{var}(\bar{z}_M^*)$ by using the dual solution:

$$\text{var}(\bar{z}_M^*) = \sum_{\ell=1}^L \rho^{\ell^2} \text{var}(\tilde{g}^\ell) = \sum_{\ell=1}^L \rho^{\ell^2} \sigma_z^2(\hat{X}^\ell). \quad (41)$$

As the lower bound means increase monotonically with the number of iterations, we obtain the largest lower bound by $\tilde{L}B^L = \bar{z}_M^{*L}$ and $\tilde{L}B^L := N(\bar{L}B^L, \text{var}(\tilde{L}B^L))$.

4.3 Stopping Rule

The analogy to the deterministic Benders decomposition algorithm we stop, if the upper and lower bound are sufficiently close. In the case of probabilistic bounds, the algorithm has to be stopped, if the upper and lower bound are indistinguishable in distribution. Based on the idea of George B. Dantzig, we check this condition by using students-t-test to determine if $s' > 0$ with 95% probability, where

$$s' = \bar{U}B^L - \bar{L}B^L + TOL \quad (42)$$

and TOL being a given tolerance.

The employment of students-t-test requires independency of the upper and lower bound distributions. As independency is not ensured in the first place as an upper bound and a binding cut in the master problem could be obtained from the same set of samples, we obtain independency by resampling the lowest upper bound before employing students-t-test. The $\hat{X}$ corresponding to the lowest upper bound and the corresponding importance distribution have to be stored. If upper and lower bounds are close to each other, which is checked by using students-t-test without fulfilling the independence requirement, we use new samples to compute an independent upper bound. Now we check if $s' > 0$ by students-t-test.

4.4 Confidence Interval

After passing the students-t-test in the last iteration, which means that the upper and lower bound means are indistinguishable, we obtain the optimum solution $\hat{X}^L, \hat{\theta}$ from the master problem. We derive from the distributions $\tilde{L}B^L$ and $\tilde{U}B^L$ a 95% confidence interval: On the left side by using the lower bound distribution and on the right side by using the upper bound distribution. We define

$$\begin{aligned} C_{\text{left}} &= 1.96\sqrt{\text{var}(\tilde{L}B^L)} \\ C_{\text{right}} &= 1.96\sqrt{\text{var}(\tilde{U}B^L)} \end{aligned} \quad (43)$$

and obtain the confidence interval

$$\tilde{L}B - C_{\text{left}} \leq Z^* \leq \tilde{U}B + C_{\text{right}} \quad (44)$$

for the final solution Z^* .

If $(C_{\text{left}} + C_{\text{right}})/\tilde{L}B^L \leq C_{tol}$, where C_{tol} is a predefined quality criteria for the confidence interval, the obtained solution is satisfying. Otherwise the sample size has to be increased and the problem has to be solved again with the increased sample size.

4.5 Improvement of the Solution

Suppose the solution with a certain sample size was not satisfying. Instead of starting from the beginning with an increased sample size we want to use the information, which we have already collected. To do this, we look for the binding cuts in the final solution, increase the sample size and recompute the binding cuts at the same $\hat{X}^\ell$, they were originally computed. This of course means that one has to store the values of $\hat{X}^\ell$ and the associated importance distributions or recompute the latter. The enlarged sample size leads to smaller variances of the binding cuts and eventually to a smaller confidence interval of the final solution. Berry-Esséen, e.g. Hall (1979), give upper bounds on the rates of convergence in the central limit theorem. Solving the master problem again with the improved binding cuts will in general not result in an intersection of the lower and upper bound. Therefore some more iterations are necessary to obtain the optimal solution according to the increased sample size. This improvement procedure could be employed iteratively until a satisfying solution is obtained. It is a possible way to improve a non satisfying solution. It may not be very efficient and there may be better ways to do so. In general we choose a sample size such that the obtained confidence interval is satisfying. We can state now the algorithm as follows:

4.6 The Algorithm

Step 0 Initialize:

$$\ell = 0, \bar{U}B^0 = \infty.$$

Step 1 Solve the relaxed master problem and obtain a lower bound:

$$\bar{L}B^\ell = c\hat{X} + \hat{\theta}^\ell.$$

Step 2 $\ell = \ell + 1$

Solve subproblems and obtain an upper bound:

$$\bar{U}B^\ell = \min\{\bar{U}B^{\ell-1}, c\hat{X}^\ell + \bar{z}(\hat{X}^\ell)\}, \text{ compute and add a cut to the master}$$

problem, using Monte Carlo (importance) sampling.

Step 3 Solve the master problem and obtain a lower bound:

$$\bar{L}B^l = c\bar{X}^l + \bar{\theta}^l.$$

Step 4 $s = \bar{U}B^l - \bar{L}B^l + TOL$

Check if $s > 0$ using students-t-test.

Step 5 Compute confidence interval and obtain a solution: $Z^*, \hat{X}, \hat{\theta}$. Stop.

Improvement of the solution:

Step 6 If $(C_{\text{left}} + C_{\text{right}})/\bar{L}B \leq C_{tol}$, stop,
otherwise got to Step 7

Step 7 Increase sample size and initialize $\bar{U}B^0 = \infty$.

Step 8 Recompute binding cuts.

$$\text{Upper bound: } \bar{U}B^l = \min\{\bar{U}B^{l-1}, C\hat{X} + \bar{z}(\hat{X}^l)\}.$$

Step 9 Go to step 3

5. Numerical Results

The method has been implemented. The Fortran code for solving general large-scale two-stage stochastic linear problems with recourse using Benders decomposition and importance sampling uses MINOS (Murtagh and Saunders (1983)), which has been adapted for this purpose, as a subroutine for solving the linear programs of the master-problem and the sub-problems. Alternatively the code can also use a modified version of Tomlin's (1973) LPM1 code of the revised simplex method as a subroutine. Versions of the code are installed on several computers, like on the IBM-3090, the Microvax-workstation, and on Personal Computers. All following test results are computed on a Toshiba laptop personal computer T5200. First we present an illustrative example, a toy problem of expansion planning of power systems which we discuss in detail. Then we derive numerical results from other small test problems. Eventually we demonstrate the solution of large-scale test problems

with numerous stochastic parameters.

The illustrative example, test problem APL1P is a model of a simple power network with one demand region. There are two generators with different investment and operating costs, and the demand is given by a load duration curve with three load levels: base, medium and peak. We index the generators with $j = 1, 2$, and the demands with $i = 1, 2, 3$. The variables, x_j , $j = 1, 2$, denote the capacities which can be built and operated to meet demands d_i , $i = 1, 2, 3$. The variable y_{ij} denotes the operating level for generator j in load level i with operating cost f_{ij} . The variable y_{is} defines the unserved demand in load level i which can be purchased with penalty cost $f_{is} > f_{ij}$. The subscript s is not an index, but denotes only an unserved demand variable. The per-unit cost to build generator j is c_j . Finally, the model is formulated with complete-recourse, which means that at any given choice of x demand is satisfied for all outcomes. In this model, building new generators competes with purchasing unserved demand through the cost function, yet there is a minimum capacity b_i which has to be built for each load level. The availabilities of the two generators, β_j , $j = 1, 2$, and the demands in each load level, d_i , $i = 1, 2, 3$, are uncertain. Generator one has four possibilities, while generator two has five, and each demand has four. All of the data values are given in Table 1 and the problem can be formulated as follows:

$$\begin{array}{ll}
\min_{s/t} \sum_{j=1}^2 c_j x_j + E\{\sum_{j=1}^2 \sum_{i=1}^3 f_{ij} y_{ij}^\omega + \sum_{i=1}^3 f_{is} y_{is}^\omega\} & \geq b_j \quad j = 1, 2 \\
x_j & \\
-\alpha_j^\omega x_j + \sum_{i=1}^3 y_{ij}^\omega & \leq 0, \quad j = 1, 2 \\
\sum_{j=1}^2 y_{ij}^\omega + y_{is}^\omega & \geq d_i^\omega, \quad i = 1, 2, 3 \\
x_j, & y_{ij}^\omega, \quad y_{is}^\omega \quad j = 1, 2, \\
& i = 1, 2, 3.
\end{array} \tag{45}$$

We will take $\omega \in \Omega$ when solving the universe problem and $\omega \in S$ when solving a problem with sampling.

Generator Capacity Costs ($10^6\$/(\text{MW}, a)$) $c_1 = 0.4, \quad c_2 = 0.25$					
Generator Operating Costs ($10^6\$/\text{MW}, a$) $f_{11} = 0.43 \quad f_{21} = 0.87$ $f_{12} = 0.20 \quad f_{22} = 0.40$ $f_{13} = 0.05 \quad f_{23} = 0.10$					
Unserved Demand Penalties ($10^6\$/\text{MW}, a$) $f_{1s} = f_{2s} = f_{3s} = 1.0$					
Minimum Generator Capacities (MW) $b_1 = b_2 = 1000$					
Demands (MW)					
#	1	2	3	4	
Outcome	900	1000	1100	1200	
Probability	0.15	0.45	0.25	0.15	
Availabilities of Generators					
Generator 1 (β_1)					
#	1	2	3	4	
Outcome	1.0	0.9	0.5	0.1	
Probability	0.2	0.3	0.4	0.1	
Generator 2 (β_2)					
#	1	2	3	4	5
Outcome	1.0	0.9	0.7	0.1	0.0
Probability	0.1	0.2	0.5	0.1	0.1

Table 1: APL1P, test problem data.

The number of possible demands and availabilities results in $4 * 5 * 4^3 = 1280$ possible outcomes in Ω , and thus 1280 subproblems have to be solved in each iteration of Benders decomposition for the universe case. We compare the universe solution with solutions gained by the importance sampling algorithm. Table 2 shows the results in the case of 20 samples out of the possible 1280 combinations and without an improvement phase. 100 replications of the same experiment with different seeds were run to get statistical information about the accuracy of the solution and the estimated confidence interval. The mean over the 100 replications of the objective function value (total costs) differs from the universe solution by 0.3%. From the distribution of the optimum objective function value derived from the 100 replications of the experiment a 95% confidence interval is computed: plus minus 2.1%. In each replication a 95% confidence interval of the solution is estimated. The mean over all replications of the estimated confidence interval is on the left side 1.5% and on the right side 1.9%. In the worst case an objective function value of 26233.9 was computed. This is about 6.4% off the correct answer. The estimated 95% confidence interval in this case did not cover the correct answer. The coverage rate of 90% expresses that in 90% of the 100 replications the correct answer of the universe solution is covered by the estimated confidence interval. This shows that if using a sample size of 20 we are slightly underestimating the confidence interval: if the computation of the 95% confidence interval was exact, we would expect a coverage rate of 95%. The reason for the underestimation of the 95% confidence interval in the case of sample size 20 lies in the underlying assumptions of the estimation method, e.g. constant error distribution along a cut, same basis for all outcomes of the random right hand sides of the cuts. Especially the latter assumption is only true, if the variances are small. A larger sample size reduces the variances and we expect a better coverage rate of the 95% confidence interval. The bias and the confidence interval of the optimum strategies (the loads x to be installed) are larger

than those of the optimum objective function value. The objective function near the optimal solution appears to be flat: several different strategies lead close to the optimum costs. Confidence intervals of about 57% and 52% are computed. In the above example a sample size of 20 was chosen. Note that additional computational effort is also needed to obtain the importance distribution, e.g. 17 subproblems have to be solved in each iteration to obtain the marginal costs M_i . Compared to the universe solution the method e.g. achieves with about 2.9% of computational effort a solution which is with 95% confidence within an interval of plus minus 2.1% of the correct answer. Importance sampling seems to be a promising approach to solving stochastic linear programs. Table 3 represents the results when using 200 samples: One can see decreasing bias, decreasing confidence intervals and improving estimations of the confidence interval with increased sample size. The coverage of the 95% confidence interval, computed by 100 replications of the experiment with different seeds, is now 95%.

We investigate the performance of the algorithm on two other examples which are small enough to compute the universe solution. PGP2, derived from Louvaux (1988), is a power generation planning model used to determine the capacities of various types of equipment required to ensure that consumer demand is met. The demands in 3 demand regions are stochastic and represented by discrete random variables with 9, 9 and 8 outcomes. CEP1 is a capacity planning model for a manufacturing plant in which several parts are produced on several machines. If the demand for the parts exceeds the production capability the residual parts are purchased from external sources at a price much higher than the production costs to meet the demand. There are 3 stochastic parameters (demands for parts), with discrete and uniform distribution with 10 outcomes each. The formulations and data for CEP1 and PGP2 may be found in Higle, Sen and Yakowitz (1990).

In the case of PGP2 we obtained very accurate results using a sample size of

50. By computing 100 replications of the experiment the mean of the objective function values differs 0.1% from the correct answer. The 95% confidence interval of the objective function value, computed by the 100 replications of the experiment is $\pm 0.76\%$, the mean of the confidence intervals estimated in each replication is on the left side 0.62% and on the right side 0.9%. In 98% the correct solution is covered by the 95% confidence interval. In the worst case the solution differed by 0.77% from the correct answer and was not covered by the 95% confidence interval.

In the case of CEP1 a higher sample size is needed to obtain accurate results. The estimation of the second stage costs appears to be more difficult. The reason lies in the fact that the (penalty) costs of buying parts from external sources are much higher than the costs of production. For this problem the additive approximation function is not a very good approximation to the true cost function as it does not cover the very high costs in scenarios where all 3 demands are high. The estimated confidence interval seems to be large, we computed 4.65% on the left side and 4.62% on the right side (mean over 100 replications of the experiment). The estimations of the confidence interval are accurate as indicated by the coverage rate of 95% of the correct answer by the 95% confidence interval. In the worst case a difference of 8.07% of the objective function value to the correct answer is computed. The worst case solution is not covered by the estimated confidence interval. In this examples it is easier to compute the value of the first stage variables than to estimate the second stage costs. In most cases the correct answer of the first stage variables was obtained. We have developed methods which adaptively improve the approximation function if sample information shows that the variance of the estimation is too large. The discussion of the adaptive approach is not subject of this paper. Table 4 and Table 5 represent the computational results of PGP2 and CEP1 and show the sizes of the test problems.

In the following we report on the solution of some large test problems with

several stochastic parameters, which are too big to be solved by computing the universe solution.

WRPM is a prototype multi-area capacity expansion planning problem for the western USA and Canada. The model is detailed covering 6 regions, 3 demand blocks, 2 seasons, and several kinds of generation and transmission technologies. The objective is to determine optimum discounted least cost levels of generation and transmission facilities for each region of the system over time. The model minimizes the total discounted costs of supplying electricity (investment and operating costs) to meet the exogenously given demand subject to expansion and operating constraints. A description of the model can be found in Dantzig et al. (1989) and Avriel, Dantzig and Glynn (1989). In the stochastic version of the model the availabilities of generators and transmission lines and demands are subject to uncertainty. There are 13 stochastic parameters per time period (8 stochastic availabilities of generators and transmission lines and 5 uncertain demands) with discrete distributions with 3 or 4 outcomes. The operating sub-problems in each period are stochastically independent. The test problem WRPM1 covers a time horizon of 1 future period, WRPM2 covers 2 future periods. There are differences in the parameters between WRPM1 and WRPM2. Note that in the deterministic equivalent formulation the problem would have more than 1.5 billion (WRPM1) and more than 3 billion (WRPM2) equations.

FI2 is a portfolio management test problem, formulated as a network problem. It is a modified version of test problems found in Mulvey and Vladimirov (1989). The problem is to select a portfolio which maximizes expected returns in future periods taking into account the possibility of revising the portfolio in each period. There are also transaction costs and bounds on the holdings and turnovers. The test problem FI2 covers a planning horizon of two future periods. The returns of the stocks in the two future periods are stochastic parameters. The problem is

formulated as a 2-stage problem. Rather than solving the problem by looking at a certain number of preselected scenarios (18 to 72 in case of Mulvey and Vladimirou) we instead assumed the returns of the stocks in the future periods being independent random parameters, discretely distributed with 3 outcomes each. As there are 13 stocks with uncertain returns, the problem has 26 stochastic parameters. The universe number of scenarios ($2.5 \cdot 10^{12}$) is very large, so that the deterministic equivalent formulation of the problem has more than 10^{14} rows. The stochastic parameters appear in the B-matrix as well as in the D-matrix.

Computational results of the large-scale test problems are represented in Table 6. Besides the solution of the stochastic problems Table 6 also shows the results from solving the expected value problem. In this case the stochastic parameters are substituted by their expectations to obtain a deterministic problem. The expected value solution is then used as a starting point for the stochastic solution. We also report on the estimated expected costs of the expected value solution. These are the total expected cost which would occur if the expected value solution is implemented in a stochastic environment. The objective function value of the true stochastic solution has to lie between the objective function value of the expected value solution and the expected costs of the expected value solution.

In case of WRPM1 and WRPM2 we chose a sample size of 100. The estimate of the objective function value of the stochastic solution (289644.2 in case of WRPM1 and 143109.2 in case of WRPM3) turns out to be amazingly accurate. The 95% confidence interval is computed as 0.0913% on the left side and 0.063% on the right side (WRPM1) and 0.0962% on the left side and 0.1212% on the right side (WRPM2). Thus the objective function value of the stochastic solution lies with 95% probability within $289379.7 \leq z^* \leq 289826.0$ (WRPM1) and $142971.5 \leq z^* \leq 143282.6$ (WRPM2). In both cases the expected costs of the expected value solution and the expected costs of the stochastic solution differ significantly. The solution time

on a Toshiba T5200 laptop PC with 80387 mathematic coprocessor was 75 minutes (WRPM1) and 187 minutes (WRPM2). During this time about 7500 (WRPM1) and 15700 (WRPM2) subproblems (linear programs of the size of 302 rows and 289 columns) get solved.

A sample size of 200 has been chosen for solving test problem FI2. The problem gets solved in only 4 iterations. The objective function value of the stochastic solution is computed as 1.1695 with a 95% confidence interval of 0.454% on the left side 0.371% on the right side. Thus with 95% probability the optimal solution lies between $1.164 \leq z^* \leq 1.174$. The estimated expected costs of the expected value solution (1.172) lie within the 95% confidence interval of the costs of the stochastic solution, however also in this case expected costs of the expected value solution and expected costs of the stochastic solution differ significantly.

6. Conclusion

We have discussed a promising approach to solving two stage stochastic linear programs with recourse and obtained first numerical results: employing importance sampling within the Benders decomposition algorithm we got very accurate solutions to the test problems with only a small sample size. The technique enables us to solve large-scale problems with a large number of stochastic parameters on a laptop computer. The test problems solved so far include up to 26 stochastic parameters and the sub-problems have a size of several hundred rows and columns. In the deterministic equivalent formulation the problems would have more than several billions of equations. The small confidence intervals of the solutions indicate that an extension to even more stochastic parameters is possible. The analysis in this paper concentrated on discrete distributions. The method, however, can be easily extended to continuous distributions. Current research concentrates on testing the technique on large-scale problems of different areas with large numbers of stochastic

parameters. We investigate possibilities to adaptively improve the approximation function, if it turns out that the error of the estimate exceeds a predefined level. If some problems require a much higher sample size the use parallel processors will enable us to quickly solve large numbers of samples to obtain low variances of the estimations. A parallel implementation of the method is in preparation.

Table 2: Model APL1P, 20 samples (100 replications of the experiment)

	correct	mean	95% conf %	bias %
#univ	1280			
#iter		7.6		
G1	1800.0	1666/5	57.0	- 7.4
G2	1571.4	1732.5	52.5	10.2
theta	13513.7	13729.4	21.3	1.6
obj	24642.3	24726.7	2.1	0.3
est. conf (%)	left	1.5		
est. conf (%)	right	1.9		
coverage		0.90		

Table 3: Model APL1P, 200 samples (100 replications of the experiment)

	correct	mean	95% conf %	bias %
#univ	1280			
#iter		7.9		
G1	1800.0	1728.7	31.5	- 4.0
G2	1571.4	1681.7	29.2	7.0
theta	13513.7	13554.7	12.2	0.3
obj	24642.3	24673.8	0.4	0.1
est. conf (%)	left	0.4		
est. conf (%)	right	0.7		
coverage		0.95		

Table 4: Model PGP2, 50 samples (100 replications of the experiment)

	correct	mean	95% conf %	bias %
#univ	648			
#iter		9.1		
obj	392.2	392.5	0.76	0.1
est. conf (%)	left	0.62		
est. conf (%)	right	0.9		
coverage		0.98		
comp. time (min)		0.28		
Problem Size				
Master:	rows	3		
	columns	7		
	nonzeros	16		
Sub:	rows	8		
	columns	16		
	nonzeros	52		

Table 5: Model CEP1, 200 samples (100 replications of the experiment)

	correct	mean	95% conf %	bias %
#univ	1000			
#iter		6.4		
obj	57790.7	58832.7	4.63	1.8
est. conf (%)	left	4.65		
est. conf (%)	right	4.62		
coverage		0.95		
comp. time (min)		0.28		
Problem Size:				
Marks:	rows	12		
	columns	10		
	nonzeros	36		
Sub:	rows	9		
	columns	16		
	nonzeros	53		

Table 6: Large test problems: computational results

	WRPM1	WRPM2	FI2
# iter stoch. (exp. val.)	139 (82)	131 (83)	4 (2)
sample size	100	100	200
exp. val. solution obj	286323.2	140041.0	1.0766
exp. val. solution, exp. cost	295473.7	147227.3	1.172
stochastic solution	289644.2	143109.2	1.169
est. conf. left %	0.0913	0.0962	0.454
conf. right %	0.063	0.1212	0.371
solution time (min)	75	187	2
Problem Size			
Master rows	44	86	48
columns	76	151	33
nonzeros	153	334	130
Sub rows	302	302	61
columns	289	289	45
nonzeros	866	866	194
# univ. scenarios	5038848	10077696	$2.5 \cdot 10^{12}$

Acknowledgements

Planning under uncertainty and solving stochastic problems by combining decomposition techniques, sampling techniques and parallel processors is a theme composed by Professor George B. Dantzig. The research leading to this paper was conducted while the author was a visiting scholar at the Department of Operations Research at Stanford University. The author wants to thank Professor George B. Dantzig for his outstanding personal and professional support during this visit. The author also wants to thank Peter Glynn and John Stone for valuable discussions on this topic. The author is grateful to James K. Ho, R.P. Sundarraj and Kingsley Gnanendran for their help in adapting the LPM1 optimizer for our purpose, to Manuel Nunez for his assistance in adapting the MINOS software and to Alamuru Krishna who assisted in preparing some of the test problems. The author is grateful to Mordecai Avriel, Julie Hagle, Suvrajeet Sen and anonymous referees for helpful suggestions concerning previous versions of this paper.

Literature

- Avriel, M., Dantzig, G.B. and Glynn, P.W. (1989): Decomposition and Parallel Processing Techniques for Large Scale Electric Power System Planning Under Uncertainty, *Proceedings of the Workshop on Resource Planning Under Uncertainty*, Stanford University.
- Benders, J.F. (1962): Partitioning Procedures for Solving Mixed- Variable Programming Problems, *Numerische Mathematic 4*, 238-252.
- Birge, J.R. (1985): Decomposition and Partitioning Methods for Multi-Stage Stochastic Linear Programming, *Operations Research 33*, 989-1007.
- Birge, J.R. and Wallace, S.W. (1988): A Separable Piecewise Linear Upper Bound for Stochastic Linear Programs, *SIAM J. Control and Optimization*, 26, 3.
- Birge, J.R. and Wets, R.J. (1986): Designing Approximation Schemes for Stochastic Optimization Problems, in Particular For Stochastic Programs with Recourse, *Math. Progr. Study 27*, 54-102.
- Dantzig, G.B. (1955): Linear Programming under Uncertainty, *Management Science 1*, 197-206.
- Dantzig, G.B. and Glynn, P.W. (1990): Parallel Processors for Planning Under Uncertainty, *Annals of Operations Research 22*, 1-21.
- Dantzig, G.B., Glynn, P.W., Avriel, M., Stone, J., Entriken, R., Nakayama, M. (1989): Decomposition Techniques for Multiarea Generation and Transmission Planning under Uncertainty, EPRI report 2940-1
- Davis, P.J. and Rabinowitz, P. (1984): Methods of Numerical Integration, Academic Press, London.
- Ermoliev, Y. (1983): Stochastic Quasi-gradient Methods and Their Applications to Systems Optimization, *Stochastics 9*, 1- 36.
- Ermoliev, Y. and R.J. Wets Eds. (1988): Numerical Techniques for Stochastic Optimization, Springer Verlag.
- Frauentorfer K. (1988): Solving SLP Recourse Problems with Arbitrary Multivariate Distributions – The Dependent Case, *Mathematics of Operations Research*, Vol. 13, No. 3, 377-394.
- Frauentorfer K. and Kall, P. (1988): Solving SLP Recourse Problems with Arbitrary Multivariate Distributions – The Independent Case, *Problems of Control and Information Theory*, Vol 17 (4), 177-205.
- Geoffrion, A.M. (1974): Elements of Large-Scale Mathematical Programming, *Management Science 16*, No 11.

- Glynn, P.W. and Iglehart, D.L. (1989): Importance Sampling for Stochastics Simulation, *Management Science*, Vol 35, 1967-1992.
- Hammersly, J.M. and Handscomb (1964): Monte Carlo Methods, Mathuen, London.
- Hall, P. (1982): Rates of Convergence in the Central Limit Theorem, *Research Notes in Mathematics* 62.
- Higle, J.L. and Sen, S. (1989): Stochastic Decomposition: An Algorithm for Two Stage Linear Programs with Recourse, to appear in *Mathematics of Operations Research*.
- Higle, J.L., S. Sen, D. Yakowitz (1990): Finite Master Programs in Stochastic Decomposition, Technical Report, Dept. of Systems and Industrial Engineering, The University of Arizona, Tucson, AZ 85721.
- Kall, P. (1979): Computational Methods for Two Stage Stochastic Linear Programming Problems, *Z. angew. Math. Phys.* 30, 261-271.
- Kall, P., Ruszczyński, A., Frauendorfer, K. (1988): Approximation Techniques in Stochastic Programming, in: Ermoliev, Y. and R.J. Wets (Eds.): Numerical Techniques for Stochastic Optimization.
- Louveaux, F.V. and Smeers, Y. (1988): Optimal Investment for Electricity Generation: A Stochastic Model and a Test Problem, in: Ermoliev, Y. and R.J. Wets (Eds.): Numerical Techniques for Stochastic Optimization.
- Mulvey, J.M. and H. Vladimirou (1989): Specifying the Progressive Hedging Algorithm to Stochastic Generalized Networks, Statistics and Operations Research Series Report, SOR-89-9, Princeton University, Princeton, New Jersey 08544.
- Murtagh, B.A. and Saunders, M.A. (1983): MINOS User's Guide, SOL 82-20, Department of Operations Research, Stanford University, Stanford CA 94305.
- Nazareth, L. and Wets, R.J.: Algorithms for Stochastic Programs: The Case of Nonstochastic Tenders, *Mathematical Programming Study* 28, 48-62.
- Pereira, M.V., Pinto, L.M.V.G., Oliveira, G.C. and Cunha, S.H.F. (1989): A Technique for Solving LP-Problems with Stochastic Right-Hand Sides, CEPEL, Centro del Pesquisas de Energia Electria, Rio de Janeiro, Brazil.
- Rockafellar, R.T. and Wets, R.J. (1989): Scenario and Policy Aggregation in Optimization under Uncertainty, *Mathematics of Operations Research* (to appear).
- Ruszczyński, A. (1986): A Regularized Decomposition method for Minimizing a Sum of Polyhedral Functions, *Mathematical Programming* 35, 309-333.
- Van Slyke and Wets, R.J. (1969): L-Shaped Linear Programs with Applications to Optimal Control and Stochastic Programming, *SIAM Journal of Applied Mathematics*, Vol 17, 638-663.

- Wets R.J. (1984): Programming under Uncertainty: The Equivalent Convex Program, *SIAM J. on Appl. Math.* 14, 89-105
- Tomlin, J. (1973): "LPM1 User's Guide", Unpublished Manuscript, Systems Optimization Laboratory, Stanford University.

REPORT DOCUMENTATION PAGE

Form Approved
OMB No 0704-0188

Public reporting burden for this collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington, VA 22202-4302, and to the Office of Management and Budget, Paperwork Reduction Project (0704-0188) Washington, DC 20503

1. AGENCY USE ONLY (Leave blank)		2. REPORT DATE March 1991		3. REPORT TYPE AND DATES COVERED Technical Report	
4. TITLE AND SUBTITLE Monte Carlo (Importance) Sampling within a Benders Decomposition Algorithm for Stochastic Linear Programs Extended Version: Including Results of Large-Scale Problems				5. FUNDING NUMBERS N00014-89-J-1659	
6. AUTHOR(S) Gerd Infanger					
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Department of Operations Research - SOL Stanford University Stanford, CA 94305-4022				8. PERFORMING ORGANIZATION REPORT NUMBER 1111MA	
9. SPONSORING MONITORING AGENCY NAME(S) AND ADDRESS(ES) Office of Naval Research - Department of the Navy 800 N. Quincy Street Arlington, VA 22217				10. SPONSORING / MONITORING AGENCY REPORT NUMBER SOL 91-6	
11a. DISTRIBUTION AVAILABILITY STATEMENT UNLIMITED				11b. DISTRIBUTION CODE UL	
12. ABSTRACT (Maximum 200 words) The paper focuses on Benders decomposition techniques and Monte Carlo sampling (importance sampling) for solving two-stage stochastic linear programs with recourse, a method first introduced by George B. Dantzig and Peter Glynn (1990). The algorithm is discussed and further developed. The paper gives a complete presentation of the method as it is currently implemented. Numerical results from test problems of different areas are presented. Using small test problems we compare the solutions obtained by the algorithm with the universe solutions. We present the solution of large-scale problems with numerous stochastic parameters which in the deterministic equivalent formulation would have billions of constraints. The problem concern expansion planning of electric utilities with uncertainty in the availabilities of generators and transmission lines and portfolio management with uncertainty in the future returns.					
14. SUBJECT TERMS large-scale				15. NUMBER OF PAGES 39 pp.	
				16. PRICE CODE	
17. SECURITY CLASSIFICATION OF REPORT UNCLASSIFIED	18. SECURITY CLASSIFICATION OF THIS PAGE	19. SECURITY CLASSIFICATION OF ABSTRACT	20. LIMITATION OF ABSTRACT SAR		